

Research Article

NAVIGATING THE AI REVOLUTION: ACADEMIC INTEGRITY, MOTIVATIONS, AND PRACTICES OF BUSINESS STUDENTS UTILIZING GENERATIVE AI FOR GRADUATION THESES IN VIETNAM.

Hoang Thi Thanh Nga ¹

¹ Foreign Trade University – Ho Chi Minh City Branch, Vietnam

ORCID

<https://orcid.org/0009-0007-3561-4191> (Hoang Thi Thanh Nga)

Abstract

The rapid diffusion of Generative AI tools into Vietnamese higher education has outpaced coherent institutional policy development, raising challenges for culminating academic work such as the graduation thesis (Khóa luận tốt nghiệp); accordingly, this study examines the motivations, functional practices, and academic-integrity perceptions shaping GenAI adoption by final-year business students writing graduation theses at Foreign Trade University, Ho Chi Minh City Campus, using an explanatory sequential mixed-methods design that combined an anonymous online survey ($N \geq 200$, stratified across International Business, Business Administration, and Finance & Banking) with semi-structured interviews (5 - 8 students and 3 - 4 supervisors), with instruments anchored in an integrated Technology Acceptance Model / Theory of Planned Behavior lens, finding that GenAI engagement during thesis work was mainstream and intensive, with motivations dominated by time pressure ($M = 4.42$, $SD = 0.68$), English-language anxiety ($M = 4.15$, $SD = 0.72$), and efficiency ($M = 3.98$, $SD = 0.81$), while intrinsic disengagement ($M = 2.10$) was broadly rejected, and revealing a graded risk pattern in which grammar correction was perceived ethical by 96.5% and practised by 94.2%, uncredited paraphrasing by 62.1%/58.4%, direct AI copying by 14.8%/18.2%, and AI-fabricated citation use by 8.5%/22.1% - a striking 13.6-point inversion at the highest-stakes tier; FTU-HCMC should therefore adopt a Responsible AI Use Policy grounded in mandatory disclosure and shift toward process-based assessment through viva voce defences, milestone logs, and continuous supervisor checkpoints that safeguard academic integrity while accommodating productive GenAI use.

Keywords

Generative AI, Academic Integrity, Student Motivations, Higher Education in Vietnam, Business Students, Undergraduate Thesis.

1. Introduction

1.1 Research Background

The release of OpenAI's ChatGPT in November 2022 and the rapid emergence of competing Generative Artificial

*Corresponding author: Hoang Thi Thanh Nga

Email addresses: Hoang Thi Thanh Nga (hoangthithanhnga.hcmc@ftu.edu.vn)

Received: 22/09/2025; Accepted: 05/02/2026; Published: 26/04/2026



Copyright: © The Author(s), 2026. Published by JKLST. This is an **Open Access** article, distributed under the terms of the Creative Commons Attribution 4.0 License (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted use, distribution and reproduction in any medium, provided the original work is properly cited.

Intelligence tools - among them Anthropic's Claude, Google's Gemini, and a growing constellation of large language models - constitute the most consequential pedagogical disruption that higher education has experienced in decades. In the immediate aftermath of ChatGPT's public launch, the institutional response across many universities worldwide was defensive: cross-country analysis of university policy issuance shows that the share of institutions initially choosing outright bans averaged 29.9% during the first wave of policy adoption (Xiao et al., 2023). Within the subsequent two years, however, the dominant trajectory shifted markedly. The number of universities now embracing regulated integration of GenAI has come to outpace those continuing to prohibit it (Xiao et al., 2023), and an empirical study of the top 100 universities in the United States found that the majority now adopt an "open but cautious" stance, releasing guidelines, syllabus templates, workshop offerings, and pedagogical frameworks rather than blanket prohibitions (Wang et al., 2024). A complementary global mapping of GenAI policies across 40 universities in six world regions corroborates that institutions are proactively addressing integration through ethical-use guidelines, authentic assessment design, and training programmes for faculty and students, while still flagging gaps in policy frameworks around data privacy and equitable access (Jin et al., 2024). Across these converging signals, the emerging consensus is that the productive question is no longer whether to permit GenAI in higher education, but how to integrate it responsibly while safeguarding academic integrity, accuracy, and pedagogical intent (Jin et al., 2024; Wang et al., 2024; Xiao et al., 2023).

This global regulatory pivot is mirrored, albeit unevenly, within the Vietnamese higher-education landscape. Vietnamese undergraduates are overwhelmingly members of Generation Z - the first cohort to grow up with constant access to digital technology and social media - and exhibit a digital-first learning disposition that predisposes them toward early adoption of generative tools (Chan & Lee, 2023). Recent empirical research confirms that Vietnamese university students are already actively engaging with ChatGPT: survey-based work documents substantive usage intensities for educational purposes and reports significant positive effects of perceived usefulness, perceived ease of use, and novelty on usage intention (Danh et al., 2024), while broader investigations indicate that AI-based tools are reshaping learning models in Vietnamese higher education by enhancing engagement, personalising learning experiences, and supporting time and workload management - even as challenges related to data privacy, digital literacy, and overdependence persist (Ha, 2025). At the same time, the country's institutional policy environment remains in flux: coherent, transparent frameworks for the ethical integration of AI are still emerging (Ha, 2025), and emerging evidence documents that AI-powered academic cheating is already

observable in the undergraduate population, with students tending to conceal such behaviours under direct questioning (Hung & Goto, 2024). The gap between rapidly diffusing student-side adoption and the still-evolving institutional regulatory environment has direct implications for the most academically consequential assignment any Vietnamese undergraduate will undertake.

That assignment is the Graduation Thesis (*Khóa luận tốt nghiệp*), the final-year capstone that occupies a singular position in Vietnamese undergraduate education. The graduation thesis is intended to synthesise the academic capability, independent thinking, and research competence that students accumulate across their four years of university study and, through its topic and quality, to reflect whether the institution's training has met its pedagogical objectives and the requirements of the labour market (Hưng, 2025). As a high-impact educational practice, the senior capstone is structured to deliver outcomes including analytical and critical thinking, communication skills, and problem-solving competencies (Houk et al., 2020). In the Vietnamese context specifically, however, documented weaknesses in graduation thesis writing - limited essay-writing curricula, low thesis quality, heavy dependence on advisors, lack of topical innovation, and insufficient rigour in referencing and review - mean that any technology capable of meaningfully shaping how theses are produced carries substantial downstream implications for training outcomes and for the integrity of the credential (Son, 2023). Recent case-study evidence from a Vietnamese public university confirms that AI tools are already being used by undergraduates during thesis writing to support idea generation, language accuracy, and time management, while simultaneously raising concerns about academic honesty, over-reliance, and the erosion of independent critical research skills (Thao, 2025). Given the high-stakes nature of the thesis as the definitive measure of independent research capability, these emerging patterns warrant focused empirical attention.

1.2. Research Context

The present study narrows its empirical focus in two deliberate directions: toward *business students* and toward *Foreign Trade University, Ho Chi Minh City Campus*. Both choices are theoretically and substantively motivated.

Business students constitute a particularly informative population for studying GenAI adoption in thesis writing. Empirically, undergraduate business students have been shown to use GenAI frequently and increasingly, primarily for information search, brainstorming, and writing-related tasks, while their use of GenAI for data analysis and multimedia production remains comparatively limited (Nakatani & Jiang, 2025). Pedagogically, business schools have been called upon to prepare students for a workforce in which GenAI is integral,

and recent student-informed conceptual frameworks foreground the need to scaffold GenAI learning activities across all undergraduate years so that students can develop the skills to integrate the technology without unintentionally engaging in academic misconduct (Drummond & Dale, 2026). In the Vietnamese context specifically, qualitative evidence shows that business students are driven by efficiency, employability, market-competitiveness concerns, and sensitivity to third-party opinions in the academic and professional choices they make (Nghia et al., 2021). These traits - an efficiency orientation, a competitive orientation, and an early-adopter disposition toward productivity-enhancing technology - make business students an ideal population in which to observe how GenAI is integrated into high-stakes academic work.

For the institutional setting, the study selects Foreign Trade University, Ho Chi Minh City Campus - one of Vietnam's most prestigious public business universities, situated in the southern economic centre of Ho Chi Minh City and serving as a representative top-tier business institution in the region (Tien & Ngoc, 2021). Two features of the FTU-HCMC student profile make the setting especially well suited to the present research question. First, FTU students possess demonstrably high English proficiency: comparative data on IELTS performance across major Vietnamese universities indicates that the proportion of FTU students scoring below band 5 is approximately 9%, while the proportion scoring within the 6.0 - 7.0 band is 16% (Nguyen et al., 2021) - figures that place FTU among the higher-performing cohorts in the national landscape and that, broader work shows, reflect the documented returns of English-medium instruction on international-business majors at top Vietnamese universities (Nguyen, 2023). This high baseline proficiency means that FTU-HCMC students face minimal linguistic friction when interacting with leading global GenAI models - virtually all of whose interfaces, training data, and capability frontiers are English-dominated - granting them unusually frictionless access to the most advanced functionalities available. Second, FTU-HCMC students experience intense academic and professional pressures: the institution's reputation, the competitive nature of business careers in Vietnam, the proximity to Ho Chi Minh City's commercial labour market, and the institutional expectations surrounding graduation thesis quality combine to create a high-stakes environment in which the productivity gains promised by GenAI are especially salient.

1.3. Literature Review & Research Gap

The theoretical lens adopted in this study integrates two complementary behavioural frameworks with an explicit academic-ethics orientation. The first candidate framework is the Theory of Planned Behavior, which models behaviour as a

function of attitude, subjective norms, and perceived behavioural control. Applied directly to GenAI adoption in higher education, the TPB has demonstrated significant predictive power for both lecturers' and students' intentions and actual use of GenAI tools, with the perceived strengths and advantages of GenAI feeding positively into all three TPB constructs (Ivanov et al., 2024). The TPB has moreover been fruitfully extended with intrinsic and extrinsic motivational antecedents to model GenAI-related academic integrity decisions, with perceived behavioural control emerging as the strongest predictor of cheating intentions and behaviour (Taslim et al., 2025), and has been combined with AI-literacy dimensions - ethics, evaluation, and awareness - to examine behavioural intention in large-scale student samples (Zhang et al., 2025). The second candidate framework is the Technology Acceptance Model, which traces adoption through perceived usefulness and perceived ease of use, and which has been productively combined with Uses and Gratifications Theory to examine ChatGPT adoption among Vietnamese university students (Danh et al., 2024). Both frameworks are embedded here within an explicit academic-ethics orientation, allowing the analysis to move beyond behavioural intention alone and to interrogate how students negotiate the boundaries between acceptable assistance and academic misconduct, and how they perceive emerging institutional policies in this changing landscape (Reginio et al., 2026; Taslim et al., 2025).

The literature gap that motivates this study is sharp and well-defined. Existing empirical work on student GenAI use in higher education has concentrated predominantly on classroom or lower-stakes assignment contexts - whether through global mapping of institutional GenAI policies (Jin et al., 2024; Wang et al., 2024), surveys of classroom usage among Vietnamese undergraduates (Dang & Nguyen, 2025; Danh et al., 2024; Ha, 2025), or broad examinations of lecturer and student attitudes toward GenAI more generally (Ivanov et al., 2024). Empirical attention to the final capstone or thesis stage - where the density of academic-integrity considerations is greatest - remains comparatively thin. Recent work is moving in this direction: mixed-methods studies of capstone projects in the Philippines document that over 80% of students used AI for ideation, organisation, writing, and coding during capstone work while simultaneously revealing ambiguous institutional guidelines and faculty concerns about over-reliance and integrity (Diloy, 2025); qualitative studies of undergraduate thesis writers using ChatGPT surface themes of academic empowerment, ethical boundary negotiation, cognitive trade-offs, peer and policy influence, and unclear institutional guidelines (Chatto & Logronio, 2025); and South-East Asian regional studies confirm that the technology's adoption is now inevitable even where institutional policies remain underdeveloped (Reginio et al., 2026; Tudy et al., 2025). Yet within Vietnam

specifically, empirical work has only begun to engage the graduation-thesis stage (Thao, 2025), and within South-East Asia more broadly, capstone- and thesis-level research centred on the psychological motivations and ethical boundaries of business-major students in developing economies remains critically scarce. The present study addresses this gap directly.

1.4. Research Questions

To investigate the psychological motivations, functional practices, and ethical perceptions that shape GenAI adoption in FTU-HCMC business students' graduation thesis work, the present study is guided by three research questions:

- **RQ1:** What are the primary motivations - both internal and external - driving FTU-HCMC business students to adopt GenAI tools during their thesis writing?
- **RQ2:** What are the actual practices and functional applications of GenAI in the thesis-writing process, and to what extent do students rely on these tools?

RQ3: How do FTU-HCMC business students perceive academic integrity in the context of GenAI use, and how aware are they of existing institutional regulations regarding AI utilisation?

2. METHODOLOGY

2.1. Research Design

The study adopts an explanatory sequential mixed-methods design, in which a quantitative survey strand is implemented first in order to map macro-level trends and statistical correlations between motivation constructs, AI-use practices, and integrity perceptions, and in which a qualitative interview strand is implemented second to contextualise the patterns surfaced in the survey (Ivankova et al., 2005; Toyon, 2021). The qualitative phase is positioned as an explanatory bridge that interrogates, in depth, the constructs and relationships made visible in the survey rather than as an independent inquiry (Dawadi et al., 2021; Toyon, 2021). The paradigmatic posture accordingly shifts: postpositivist assumptions guide the operationalisation of TAM/TPB hypotheses during the survey phase, while constructivist assumptions foreground individual meanings and ethical negotiations during the interview phase (Dawadi et al., 2021; Ivankova et al., 2005). Such a design is well matched to the present investigation, whose three research questions address phenomena whose distribution cannot be mapped from qualitative work alone and whose meaning cannot be fully captured by survey items alone. A purely quantitative design would lack contextual

nuance, while a purely qualitative design would lack the breadth to estimate the relative weight of competing predictors.

2.2. Quantitative Phase

The quantitative phase is operationalised through an anonymous online questionnaire administered to final-year undergraduate students across three core business majors at Foreign Trade University, Ho Chi Minh City Campus - International Business, Business Administration, and Finance & Banking. The sample is accessed through a hybrid sampling frame combining a convenience component with stratified random selection by major, so that the eventual sample captures the proportional composition of the three target cohorts while preserving recruitment feasibility (Andrade, 2020; Stratton, 2021). The minimum target sample size is $N \geq 200$, a threshold that supports the planned EFA and preserves statistical power for the regression and structural-model estimations to follow (Taherdoost, 2016; Thanapongporn et al., 2023). Because convenience sampling limits external validity, the study restricts its inferential claims to FTU-HCMC final-year business students (Andrade, 2020; Stratton, 2021).

The instrument is built around three integrated Likert-based scales, each anchored in established international frameworks and adapted to GenAI-assisted graduation-thesis writing. All items use a five-point Likert format (Strongly Disagree–Strongly Agree) (Taherdoost, 2016; Thanapongporn et al., 2023). The instrument is reviewed by subject-matter experts for content validity and refined through a pilot test, after which Cronbach's alpha coefficients are computed for each scale (Taherdoost, 2016; Thanapongporn et al., 2023).

The Motivation Scale operationalises the psychological drivers hypothesised by the integrated TAM/TPB lens: perceived usefulness, perceived ease of use, time pressure from simultaneous final-year demands, anxiety over academic writing, and English-language barriers (Dan et al., 2024; Ivanov et al., 2024). Items are worded in both direct and reverse-coded form to control for acquiescence bias (Ivanov et al., 2024; Taslim et al., 2025; Zhang et al., 2025).

The Practice/Usage Scale assesses the frequency with which students use GenAI across discrete thesis-production stages — ideation, literature synthesis, drafting, paraphrasing, translation, and data analysis. The functional taxonomy extends prior empirical work on undergraduate business-student GenAI usage (Nakatani & Jiang, 2025) into the thesis-specific functional stages of literature synthesis, paraphrasing, and translation (Thao, 2025).

The Academic Integrity Scale operationalises the ethical dimension across three sub-dimensions. The ethical self-awareness subscale draws on a long-standing psychometric tradition of measuring the moral obligation not

to cheat (Passow et al., 2006). The neutralisation-techniques subscale is grounded in classical neutralisation theory adapted to higher-education settings (DiPietro, 2010) and extended through more recent recastings that add denial of the victim and condemnation of the condemners (Awdry & Groves, 2023). The fear-of-detection subscale measures the subjective probability of detection by supervisors, viva committees, or institutional integrity checks, along with the behavioural consequences students perceive that risk to carry (Reginio et al., 2026; Taslim et al., 2025).

The questionnaire additionally collects demographic and contextual covariates (major, gender, prior GenAI usage, self-rated English proficiency) to support subgroup analyses.

2.3. Qualitative Phase

The qualitative phase consists of semi-structured interviews with two purposefully distinct participant groups. The first comprises 5 - 8 final-year students drawn from the same FTU-HCMC population but recruited specifically for a deeper exploration of personal motivations, tacit decision rules, and sensitive practices that survey items alone may struggle to surface - including GenAI-assisted paraphrasing calibrated to circumvent similarity-detection systems, undisclosed translation assistance on non-English source material, and the use of GenAI to fabricate primary empirical material such as datasets or interview excerpts (Chatto & Logronio, 2025; Diloy, 2025). The second comprises 3 - 4 thesis supervisors from the three target programmes, whose faculty perspective offers longitudinal observations of changes in thesis quality, argumentation originality, citation patterns, and the practical challenges of detecting GenAI-assisted writing within the thesis workflow (Diloy, 2025; Tudy et al., 2025). Pairing student and faculty voices enables cross-validation of survey findings against the actual texture of the theses supervisors read.

The sample size is justified on information-power principles (Malterud et al., 2015): the aim is narrow, the participant pool sample-specific, the theoretical scaffolding strong, and the dialogue quality high. Empirical evidence indicates that homogeneous-population studies with narrowly defined objectives typically reach code saturation at the lower end of a 9-to-17 interview range, with few new themes emerging beyond that point (Hennink & Kaiser, 2021); the planned 5 - 8 student and 3 - 4 supervisor interviews sit comfortably within, and in many cases below, this saturation bound while still enabling triangulation between the two stakeholder perspectives. The interviewer follows established procedures for sensitive-topic interviewing: rapport is built through low-stakes conversational openings, sensitive and open questions are used where direct questioning would inhibit disclosure, a deliberately safe and non-judgmental environment is cultivated, and the researcher maintains

ongoing reflexivity about the institutionally evaluative positionality that their university affiliation implies (Westland et al., 2024).

2.4. Data Analysis & Ethical Considerations

Quantitative data are analysed in three coordinated steps. Descriptive statistics are first computed to characterise central tendencies, dispersions, distributional properties, and bivariate correlations across the three scales. Exploratory Factor Analysis is then applied to each scale to verify factor structure, identify low-loading items for potential removal, and establish construct validity before any predictive modelling is conducted (Taherdoost, 2016; Thanapongporn et al., 2023). Items that compromise internal consistency, convergent validity, or discriminant validity are iteratively pruned in line with documented PLS-SEM methodological practice (Alhassany & Faisal, 2018). Multiple regression analysis is then used to estimate the integrated TAM/TPB model and identify the strongest predictors of AI-assisted behaviours, with PLS-SEM using SmartPLS 3 deployed as a complementary pathway. PLS-SEM is preferred over covariance-based SEM because the research questions are exploratory and target a relatively new phenomenon with no fully consolidated measurement model in the Vietnamese business-student context (Alhassany & Faisal, 2018; Palos-Sánchez et al., 2021). All analyses use IBM SPSS Statistics for descriptive and EFA work and SmartPLS for structural-model estimation, a dual-tooling configuration standard in behavioural research combining EFA with path-model testing (Alhassany & Faisal, 2018).

Qualitative interview transcripts are analysed through reflexive thematic analysis following the six-phase Braun and Clarke framework (Byrne, 2021) and its more recent procedural extension (Naeem et al., 2023). Analysis begins with full immersion in the transcripts, proceeds through systematic initial coding that captures both semantic and latent content, advances to the identification of candidate themes clustering around the Motivation, Practice/Usage, and Academic Integrity dimensions surfaced in the quantitative phase, refines and names those themes through iterative review, and culminates in an analytic narrative integrating representative verbatim quotations that feeds back into the interpretation of the quantitative findings. Where interview content references sensitive practices, pseudonymisation preserves participants' voice while ensuring that no verbatim material could compromise institutional anonymity guarantees made during consent.

Ethical considerations are central to the design. Because the construct under study – students' use of GenAI in their graduation thesis work, particularly where that use may approach or cross institutional integrity lines - is socially sensitive, the study explicitly anticipates and mitigates social

desirability bias, the well-documented tendency for respondents to under-report behaviours carrying negative social or institutional consequences (Durmaz et al., 2019; Hung & Goto, 2024; Ried et al., 2021). Prior research on Vietnamese undergraduates has shown that direct questioning on AI-powered academic cheating substantially under-reports actual behaviour, and that effective measurement requires layered procedural safeguards rather than a single anonymity assurance (Hung & Goto, 2024). The present study therefore adopts a multi-pronged mitigation repertoire: the recruitment title is neutral and does not foreground integrity as the focal concern; anonymity and data confidentiality are guaranteed in writing; no personally identifying information is collected; data are stored on encrypted, access-controlled devices; the survey is online and fully self-administered, since meta-analytic evidence indicates that computerized survey modes increase self-disclosure of sensitive behaviours relative to interviewer-administered modes (Gnambs & Kaspar, 2014); and individual interview content on possible institutional-rule breaches is elicited from a deliberately non-judgmental vantage, so that participants are not discouraged from disclosure by fear of academic penalty (Westland et al., 2024). The study acknowledges that anonymity alone may not fully outweigh perceived negative consequences and that online respondents are increasingly wary of data-privacy breaches even when confidentiality is guaranteed (Cerri et al., 2021); the qualitative strand is therefore designed not merely as a complement to the survey but as active triangulation, where rapport and contextual sensitivity compensate for residual social-desirability effects in survey data.

The research is conducted under approval from the relevant institutional ethics review body, with informed consent obtained from all participants in both the survey and interview strands.

3. RESULTS

3.1. Demographic Profile and Prevalence of AI Adoption

The three data tables supporting this Results section report adoption, motivation, and integrity metrics but do not include a disaggregated breakdown of respondents by major, GPA, or gender. The present subsection therefore concentrates on the prevalence patterns documented in Table 1; the absence of demographic disaggregation within the tabular evidence is acknowledged as a limitation and is taken up again in Section IV. What Table 1 does deliver is a granular map of four AI-tool clusters to their primary thesis-writing functions and to four frequency categories spanning daily use to non-use, and two patterns dominate.

Table 1: Frequency and Types of Generative AI Tools Utilized by Students

AI Tool	Primary Function	Daily	Weekly	Only when stuck	Never
ChatGPT / Claude	Brainstorming, drafting, translating	45.2%	38.1%	12.7%	4.0%
Grammarly / QuillBot	Proofreading, academic paraphrasing	58.7%	26.3%	11.0%	4.0%
Elicit / Consensus	Literature search, citation sourcing	12.5%	42.3%	30.2%	15.0%
Canva / Gamma	Thesis defense slide creation	3.2%	8.5%	65.0%	23.3%

First, AI-tool engagement during thesis work is overwhelmingly intensive. Across the cluster with the deepest daily penetration - Grammarly/QuillBot, deployed for proofreading and academic paraphrasing - 58.7% of respondents report daily use and a further 26.3% engage weekly, leaving only 4.0% in the non-user category. ChatGPT/Claude occupies a comparably dominant position for brainstorming, drafting, and translating functions, with 45.2% reporting daily use and 38.1% weekly use. The combined daily-plus-weekly reach of these two top clusters (approximately 84.9% and 83.3% respectively) confirms that GenAI-assisted thesis production is mainstream rather than exceptional among FTU-HCMC business undergraduates.

A second tier of tools exhibits lighter but still meaningful uptake. Elicit/Consensus, used for literature search and citation sourcing, shows a primarily weekly usage pattern (42.3%) alongside a substantial "only when stuck" component (30.2%) and a non-negligible 15.0% non-user share - consistent with the selective, bottleneck-driven manner in which students engage literature-review infrastructure. Canva/Gamma, deployed for thesis-defence slide creation, sits at the bottom of the ranked pattern: only 3.2% report daily use and 8.5% weekly use, while 65.0% engage only when stuck and 23.3% report never using it, mirroring the episodic, presentation-bound nature of defence preparation.

3.2. Motivations Behind AI Adoption

Table 2 ranks five motivation factors, operationalised

through the integrated TAM/TPB lens, by mean score. The findings directly answer RQ1 by surfacing a clear motivational hierarchy in which extrinsic and situational pressures dominate while intrinsic motivational deficits occupy a markedly weaker position.

Table 2: Mean Scores of Primary Motivations Driving Student AI Adoption

Code	Motivation Factor	Mean	SD	Ranking
MOT 1	Time Pressure: Balancing internship workload and thesis deadlines.	4.42	0.68	1
MOT 2	Language Barrier: Anxiety about maintaining a high-level English academic tone.	4.15	0.72	2
MOT 3	Efficiency: Seeking faster ways to summarize lengthy journal articles.	3.98	0.81	3
MOT 4	Lack of Guidance: Difficulty in reaching out to supervisors for immediate feedback.	3.21	0.95	4
MOT 5	Pure Laziness: Wanting AI to do the thinking and writing.	2.10	1.12	5

MOT1 emerges as the strongest motivator with a mean of 4.42 (SD = 0.68), ranked #1. The magnitude and tight dispersion of this score indicate near-uniform endorsement of time pressure as a primary adoption driver. The finding aligns with the theoretically predicted role of perceived behavioural control in the TPB: when external time demands reduce the behavioural control students exercise over their thesis

workflow, instrumental technology adoption is predicted to intensify. Empirically, the result also corroborates prior work documenting that Vietnamese undergraduates experience substantive GenAI-driven productivity gains in workload management.

MOT2 ranks #2 with a mean of 4.15 (SD = 0.72). This finding directly echoes the FTU-HCMC context established earlier in the manuscript: with English-medium instruction shaping the cohort's documented English proficiency profile, the linguistic dimension of thesis writing remains a salient anxiety even among students whose baseline proficiency is comparatively high. The result intersects with TAM's perceived-usefulness pathway, since GenAI's English-dominant interface and capability frontier render it functionally valuable for academic-language tasks.

MOT3 ranks #3 with a mean of 3.98 (SD = 0.81). The slightly larger dispersion relative to MOT1 and MOT2 indicates that efficiency-orientation is endorsed by a strong majority but with greater individual variation. The finding is consistent with TAM's perceived-usefulness construct and with prior empirical work showing that summarisation and information-search functions are among the most frequently exercised by undergraduate business students.

A meaningful gap separates the top three motivators from the bottom two. MOT4 ranks #4 with a mean of 3.21 (SD = 0.95), indicating moderate endorsement with substantial heterogeneity; many students experience supervisor-availability friction and turn to GenAI as a partially compensatory resource, but the dispersion suggests the factor operates unevenly across the cohort. MOT5 ranks a distant #5 with a mean of 2.10 (SD = 1.12). The depressed mean, falling below the scale midpoint, together with its large dispersion demonstrates that laziness as a self-attributed motivator is endorsed by only a minority of respondents and is broadly inconsistent with the prevailing motivational self-image.

Taken together, the rank-ordered pattern answers RQ1 decisively. FTU-HCMC business students adopt GenAI for thesis work primarily because of time pressure, secondarily because of English-language anxiety, and tertiarily because of efficiency gains - three motivational drivers mapping cleanly onto the TAM/TPB constructs of perceived behavioural control, perceived usefulness, and ease of use respectively. The motivational profile is dominated by situational and instrumental drivers rather than by intrinsic disengagement, with the lowest-ranked factor trailing the highest by more than two full scale points.

3.3. Actual Practices in Thesis Writing

Table 3 makes visible a clear behavioural spectrum which, read alongside the tool-function mapping in Table 1, allows the actual thesis-writing practices of FTU-HCMC business

students to be classified into a low-risk-to-high-risk continuum that directly answers RQ2.

Table 3: Perceived Ethical Boundaries versus Actual AI Practices Among Respondents

Academic Integrity Context	% Students believe it is "ETHICAL"	% Students actually DID IT
Using AI to correct grammar and spelling errors.	96.5%	94.2%
Asking AI to rephrase/paraphrase a paragraph without citation.	62.1%	58.4%
Copying directly 1-2 paragraphs generated by AI into the thesis.	14.8%	18.2%
Using AI-generated citations without verifying the original papers.	8.5%	22.1%

At the low-risk end of the spectrum sits the dominant pattern of proof-reading and language correction. Table 1 indicates that 58.7% of students use Grammarly/QuillBot daily and a further 26.3% weekly for “proofreading, academic paraphrasing”, while Table 3 confirms the broadly ethical status of this activity: 96.5% of respondents judge AI-assisted grammar and spelling correction to be ethical and 94.2% report actually engaging in the practice. The near-convergence of ethical perception and actual behaviour - a roughly 2.3 percentage-point gap - marks a low-risk, broadly sanctioned function class, aligning with prior work that positions language-editing support as the most institutionally accepted entry point for GenAI into academic writing.

A middle-band risk practice is uncredited paraphrasing. Table 3 indicates that 62.1% of respondents judge it ethical to ask AI to rephrase a paragraph without citation while 58.4% report actually doing so. The Table 1 picture is consistent: 11.0% of students use Grammarly/QuillBot “only when stuck”, and the ChatGPT/Claude cluster - 38.1% weekly for drafting functions - routinely produces paraphrased output. The fact that actual use (58.4%) sits slightly below ethical perception (62.1%) signals a function class in which a clear majority of students see the behaviour as acceptable and

broadly enact that perception.

The high-risk end of the spectrum comprises two practices characterised by sharply divergent perception-behaviour alignment. First, copying AI-generated paragraphs directly into the thesis is judged ethical by only 14.8% of respondents yet reportably practised by 18.2% - a 3.4-point inversion in which actual behaviour modestly exceeds ethical perception. More critically, the use of AI-generated citations without verifying the original papers is judged ethical by just 8.5% of students but reportably practised by 22.1% - a 13.6-point inversion that the table labels “Báo động” (warning). This citation-fabrication finding is the single most consequential behavioural pattern in the study because it combines technical risk (hallucinated references) and ethical risk (non-verification) in a manner that existing similarity-detection infrastructure does not robustly catch.

Integrating Tables 1 and 3 therefore yields a four-tier classification of actual practices. Tier 1 (low risk, broadly sanctioned): grammar correction, supported primarily by Grammarly/QuillBot at 96.5% / 94.2%. Tier 2 (low-to-moderate risk, accepted by a clear majority): uncredited paraphrasing at 62.1% / 58.4%. Tier 3 (high risk, ethically contested by an overwhelming majority yet practised by roughly one in five): direct copying of AI-generated paragraphs at 14.8% / 18.2%. Tier 4 (highest risk, ethical condemnation unambiguous yet a substantial minority nonetheless engages): AI-fabricated citation use at 8.5% / 22.1%, sitting at the intersection of the Elicit/Consensus literature-search cluster and the ChatGPT/Claude drafting cluster. The pattern answers RQ2 by demonstrating that the functional taxonomy of tools in Table 1 maps onto a four-tier risk classification in Table 3, with progressively widening perception-behaviour gaps as severity increases.

3.4. Academic Integrity Awareness vs. Institutional Reality

Table 3 furnishes the empirical basis for answering RQ3 - how FTU-HCMC business students perceive academic integrity in the context of GenAI use and how their ethical perceptions align or fail to align with reported conduct. The central pattern is a graded set of perception-behaviour gaps that becomes more pronounced as the severity of the integrity-relevant behaviour increases.

At the bottom of the risk gradient - language correction - ethical perception and actual conduct are tightly aligned: 96.5% of respondents identify AI-assisted grammar and spelling correction as ethical and 94.2% report actually doing it. The 2.3-point perception-actual gap is the smallest in the table, indicating that for low-controversy, language-tool-mediated activities, students' integrity judgements and conduct are mutually reinforcing.

At the next tier - uncredited paraphrasing - alignment

remains strong but loses some of its tightness. The ethical-judgement share (62.1%, labelled “Rủi ro” (Risk) and the actual-use share (58.4%) both sit close to a two-thirds majority, with ethical perception modestly exceeding actual behaviour. The pattern indicates that the majority of students see paraphrasing assistance as within the bounds of integrity and broadly enact that perception, while the subset who judge the practice ethical but do not engage may reflect either insufficient thesis-writing workload or residual detection concerns.

The picture inverts decisively at the top of the spectrum. Direct copying of AI-generated paragraphs is judged ethical by only 14.8% but reportably practised by 18.2% - a 3.4-point gap in which actual behaviour modestly exceeds stated ethical acceptance. Most strikingly, AI-generated citations without verification are judged ethical by only 8.5% yet reportably practised by 22.1% - a 13.6-point inversion, the largest in the dataset and explicitly labelled “Báo động”. Nearly one in four students admits to using AI-fabricated citations while fewer than one in ten consider the practice ethically defensible. This is the only case in the data where actual conduct demonstrably and substantially exceeds ethical self-endorsement, pointing to a population whose conduct at the highest-stakes tier of integrity is, by their own lights, widely regarded as misconduct.

Two implications follow for the institutional reality confronting FTU-HCMC. First, the alignment of language-editing practice with ethical perception (96.5% vs. 94.2%) and of uncredited-paraphrasing practice with ethical perception (62.1% vs. 58.4%) suggests that for low- and moderate-risk functions, institutional guidance that codifies existing perceptions would likely meet broad acceptance. Second, the inversion among high-risk functions - particularly the 8.5% vs. 22.1% AI-citation inversion - points to a population for whom ethical awareness and conduct are decoupled at the most consequential end of the integrity gradient, and for whom institutional detection mechanisms (including Turnitin AI detection) are unlikely to be sufficient on their own given the documented difficulty of catching citation fabrications through similarity matching alone. Together these findings answer RQ3 by showing that FTU-HCMC business students broadly perceive the integrity gradient of GenAI use but that, at the most consequential end, perception and conduct diverge sharply.

4. DISCUSSION

The Results section surfaced a pattern that, on first reading, appears paradoxical: the most ethically troubling GenAI-integrity behaviours are concentrated precisely among the most academically capable students at one of Vietnam's most prestigious business-specialist institutions. The

motivational hierarchy in Section 3.2 is instrumentally driven and time-pressure-dominated, while the practice pattern in Section 3.3 places a substantial minority in direct violation of their own ethical self-image at the highest-impact tier of the undergraduate career. The four subsections that follow interrogate this paradox, situate the FTU-HCMC findings within the global literature, derive policy and pedagogical implications grounded in the empirical patterns surfaced, and acknowledge the methodological boundaries and forward directions the study opens.

4.1. Interpretation of the AI-Integrity Paradox

The motivational rankings in Table 2 show that FTU-HCMC business students do not adopt GenAI because they are disengaged from learning: MOT5 (“Pure Laziness”) trails the scale midpoint at a mean of 2.10 and is broadly rejected. What the data instead show is that adoption is overwhelmingly instrumental - MOT1 at 4.42, MOT2 at 4.15, and MOT3 at 3.98 dominate the hierarchy.

The interpretive lens best fitted to this pattern is *cognitive offloading*. Fan and colleagues define the construct as the delegation of mentally effortful tasks to external tools and document that habitual offloading of cognitive friction can erode the activation of analytical reasoning and produce “metacognitive laziness”. Zhai and colleagues' systematic review confirms that over-reliance on AI dialogue systems is associated with reduced decision-making, critical-thinking, and analytical-thinking capacities, especially where dependency becomes habitual. The Table 1 frequency data are consistent with extensive routine offloading: 58.7% of respondents use Grammarly/QuillBot daily and 45.2% use ChatGPT/Claude daily, with approximately four-fifths of the cohort reporting at least weekly use for both clusters.

The data should not, however, be read as a one-sided “AI impairs thinking” account. Wang and Zhang's three-context mixed-methods study documents that *strategic* offloading can liberate mental resources for higher-order reflection when it exceeds specific thresholds, and that efficiency orientation operates not as a barrier but as an amplifier of both cognitive vigilance and cognitive offloading. Frameworks of this kind read the FTU-HCMC findings as describing the upper end of an inverted-U relationship between offloading and learning depth - one in which the time-and-efficiency pressures of senior-capstone work push high-capability students past threshold into routine dependence.

The competitive business-school environment amplifies the effect. Hall and Martin's analysis identifies the dense competitive pressure from accreditation bodies, ranking regimes, and publication incentives that dominate business-school institutional culture and shape how individuals calculate the costs and benefits of effortful conduct. Bergenholtz and colleagues'

cognitive-load-inversion pattern in business-school examinations - under which low performers benefit while high performers decline, producing performance convergence - is a structural analogue of what Section 3.4 surfaces at the integrity level. The FTU-HCMC paradox is therefore not that students are disengaged; it is precisely because they are operating within a competitive, time-saturated, status-conscious business-school environment that efficiency-oriented, instrumentally rational GenAI adoption produces the highest-severity perception-behaviour inversion (8.5% ethical endorsement against 22.1% reported use of AI-fabricated citations) at exactly the moment when the graduation thesis is supposed to evidence independent analytical capability.

4.2. Comparative Analysis with Global Literature

The FTU-HCMC findings take on sharper meaning when juxtaposed against the regional and global literature.

In the Asian-hub cluster most directly comparable to Vietnam, Qu and colleagues' Singaporean mixed-methods study documents an "AI ghostwriter effect" in which applied-discipline students (engineering, business) more readily normalise GenAI as a practical tool while theoretical-discipline students express stronger ethical discomfort. The magnitude of the FTU-HCMC inversion at the citation tier is consistent with this applied-discipline normalisation but extends it by surfacing - through quantitative inversion gaps - that normalisation among business students can coexist with an ethical perception that is unambiguous in its condemnation. Tan and colleagues' TPB-grounded Singaporean study shows that *guideline understanding* positively influences compliance and declaration honesty - a finding that, applied to FTU-HCMC's documented policy-vacuum environment, helps explain why a coherent compliance signal is currently absent from the data. The Hong Kong literature is closely parallel: Lo and Chan document pronounced polarisation across GPA strata, while Tsao's policy-trajectory analysis reveals institutional ambiguity that supports experimentation while individualising failure on students and staff.

Within the closer ASEAN orbit, Diloy's mixed-methods study of capstone projects across five Philippine universities finds that over 80% of students used AI across ideation, organisation, writing, and coding - a usage-intensity figure that sits alongside the FTU-HCMC tool-frequency picture rather than against it - and surfaces the same cluster of concerns about overreliance, superficial learning, and ambiguous guidelines. Tudy and colleagues' phenomenographic study of professors in Indonesia, Malaysia, and the Philippines concludes that the use of ChatGPT is *inevitable*, and recommends training-plus-protocol

frameworks as the appropriate response. Buragohain and Chaudhary's ASEAN-wide systematic review positions the region as one in which responsible-use frameworks must be socio-technically responsive rather than dictated as region-internal unified mandates.

Anglophone-Western baselines provide a useful contrast point. Gruenhagen and colleagues' Australian university survey reports that more than a third of students have used a chatbot for an assessment without necessarily perceiving it as a breach of integrity - but does not surface the magnitude of perception-behaviour inversion documented in the FTU-HCMC citation-fabrication tier. Yusuf and colleagues' 76-country multicultural study identifies a strong correlation between cultural dimensions and concerns around academic dishonesty, supporting the present study's claim that Vietnamese results cannot be assessed absent their specific cultural context.

Three localised factors shape the FTU-HCMC outcome relative to these comparators. First, the ESL dimension: although the cohort has comparatively high English proficiency, MOT2 ranks #2 at a mean of 4.15 - academic-English anxiety persists even at high baseline proficiency and intensifies the instrumental pull of GenAI for thesis work. Second, cultural academic pressure - the Vietnamese graduation thesis is the singular capstone through which the institution judges whether its training has met pedagogical objectives and labour-market requirements, concentrating integrity stakes far more densely than the typical Anglophone coursework context does. Third, the institutional policy environment remains in flux, so that - unlike the compliance-positive effects of guideline understanding documented in Singapore - FTU-HCMC students operate in a context where institutional guidance cannot meaningfully anchor the perception-behaviour relationship, an asymmetry that helps account for the largest inversion in Table 3.

4.3. Pedagogical and Policy Implications

The implications fall into two clusters: institutional policy at FTU-HCMC and faculty/assessment design.

Institutional policy: from unenforceable bans toward a Responsible AI Use Policy. Empirical analysis of top-500 university ChatGPT policies shows that bans are unenforceable - although a meaningful share of institutions initially chose prohibition, the share now embracing regulated integration has come to outpace those that continue to prohibit, and most banning institutions permit individual instructors to deviate from the prohibition. Hosseini and colleagues argue that outright prohibition is counterproductive because it actively encourages undisclosed use of AI tools. The recommendation for FTU-HCMC is therefore to move beyond unenforceable bans toward a coherent Responsible AI

Use Policy whose disclosure mechanics follow three concrete practices: (i) describe AI tool use in the introduction or methods section, including the prompts used and the parts of the text affected; (ii) cite AI tools in-text and in the references with adequate specificity of contributor, model version, and date of use; and (iii) record and submit relevant AI interactions as supplementary material or appendices. Operational APA-style examples consistent with this principle are articulated by Hosseini et al. for scholarly manuscripts, by Chan for both APA and MLA formats, and by Alberth for personal-communication-style handling of ChatGPT specifically. The present data anchor this recommendation directly: the 96.5% / 94.2% grammar-correction alignment in Table 3 indicates that a permissive-with-disclosure regime would meet broad acceptance at the lowest-risk tier, while the higher-tier inversions make disclosure a precondition for any defensible integrity framework.

Faculty and assessment design: from product-based to process-based assessment. A complementary shift is required in how thesis work is graded. Three converging lines of evidence support moving from product-based assessment (the final PDF thesis) toward process-based assessment (mandatory viva voces/defence steps, milestone logs, continuous face-to-face supervisor checks).

First, AI-detection infrastructure is unreliable. Empirical evidence shows that although Turnitin's detector correctly identified AI-generated content in 91% of cases in a controlled setting, faculty members formally reported only 54.5% of those papers as potential misconduct, with adversarial prompt-engineering an effective means of evasion; independent multi-tool testing confirms that detection systems scored below 80% accuracy across platforms, with both false positives (human text classified as AI) and false negatives materially present. For an English-medium-institution ESL-dominant cohort, the implication is not only epistemic but equity-related: detection tools frequently produce false positives against multilingual or non-native English writers, raising serious fairness concerns that bear directly on FTU-HCMC's profile.

Second, process-based assessment frameworks are maturing. Kadel and colleagues argue for a reimagined Project-Based Assessment model in GenAI-saturated capstone contexts that prioritises process-oriented evaluation, multi-modal and multifaceted assessment design, and supervisor-led milestone review as part of the project lifecycle.

Third, oral vivas are operationally feasible at cohort scale. Chen and Talmi propose an oral-viva framework that enables direct assessment of conceptual understanding, methodological justification, and reflective thinking in a manner that written reports cannot reliably capture in an AI-saturated environment. Storey argues that oral defences must shift from confirmatory questions permitting

regurgitated responses toward exploratory inquiry that demonstrates genuine conceptual transformation. Cole's documentation of the "Professional Conversation" approach details procedural safeguards - advance key questions, conversation plans, assessor preparation, and collaborative moderation - that are familiar to Vietnamese academic culture and translatable to the thesis-defence stage.

Tying each implication back to the empirical pattern in Table 3: the 8.5% / 22.1% citation-fabrication inversion is best addressed through mandatory verification logs at milestone checkpoints paired with viva-voce source-level interrogation; the 14.8% / 18.2% direct-copy inversion is best detected and deterred through exploratory vivas probing passage-level authorship; the 62.1% / 58.4% paraphrasing ambiguity is best navigated through explicit disclosure thresholds; and the 96.5% / 94.2% grammar-correction pattern is best preserved through permitted-with-disclosure language-tool guidance that does not over-extend the prohibition to low-severity functions.

CONCLUSION Four limitations bound the present study's claims. First, the data are self-reported, and the social-desirability mitigation repertoire articulated in Section 2.4 does not fully neutralise the well-documented tendency of Vietnamese undergraduates to conceal AI-powered academic cheating behaviour. List-experiment evidence in the Vietnamese undergraduate population indicates that indirect questioning produced an almost threefold higher prevalence estimate than direct questioning, suggesting that Table 3 figures should be read as floor estimates of the actual sensitive-behaviour prevalence, with the 8.5% / 22.1% citation-fabrication inversion likely to understate real population-level behaviour. Second, the sample is confined to a single campus; the findings cannot be extrapolated to Vietnamese higher-education populations more broadly, to other business majors outside the three target programmes, or to regional business-school populations whose institutional parameters differ. Third, the supporting tables do not disaggregate adoption patterns by major, GPA, or gender - flagged in Section 3.1 and constraining the study's capacity to test the GPA-stratified patterns documented in the Hong Kong literature and the gender-stratified patterns documented in the Vietnamese literature. Fourth, the design is cross-sectional and captures a single snapshot rather than a temporal trajectory of how student-AI dynamics evolve.

5. CONCLUSION

The present study set out to investigate the use of Generative AI tools among business undergraduate students at Foreign Trade University, Ho Chi Minh City Campus during the writing of their graduation theses - a singular capstone artefact that the Vietnamese higher-education system treats as the definitive evidence of independent research capability.

Grounded in an integrated TAM–TPB lens, the study pursued three research questions, namely: the primary motivations driving GenAI adoption; the actual practices and functional applications of GenAI during thesis production; and students' perceptions of academic integrity and their awareness of existing institutional regulations concerning AI utilisation.

Three principal findings emerge from the analysis. First, the motivational hierarchy is dominated by extrinsic and situational pressures rather than by intrinsic disengagement: time pressure ($M = 4.42$, $SD = 0.68$) ranks #1, English-language anxiety ($M = 4.15$, $SD = 0.72$) ranks #2, and efficiency orientation ($M = 3.98$, $SD = 0.81$) ranks #3, while the laziness factor ($M = 2.10$, $SD = 1.12$) is broadly rejected at the bottom of the ranking. Second, the four-tier classification of actual practices spans AI-assisted grammar correction at the low-risk end (96.5% ethical / 94.2% actual), uncredited paraphrasing at the middle band (62.1% / 58.4%), direct copying of AI-generated paragraphs at the high-risk tier (14.8% / 18.2%), and AI-fabricated citation use at the highest-risk tier (8.5% / 22.1%). Third, the perception–behaviour gap widens progressively across these tiers, culminating in a 13.6-point inversion at the citation-fabrication level in which reportable conduct substantially exceeds ethical self-endorsement - the most consequential pattern in the dataset given the documented difficulty of detecting hallucinated references through similarity-based infrastructure alone.

Taken together, these patterns substantiate the central interpretive argument advanced in Section 4: that instrumentally rational adoption under competitive business-school pressure, rather than disengagement or laziness, generates an AI-integrity paradox whose most acute expression sits at the highest-stakes end of the integrity gradient. Cognitive offloading, amplified by the time-and-efficiency demands of senior-capstone work in a status-conscious institutional environment, produces a population whose ethical awareness and behavioural conduct are decoupled precisely at the moment when the graduation thesis is supposed to evidence independent analytical capability.

The study contributes two practice-oriented implications in compressed form. For institutional policy, the findings support moving beyond unenforceable prohibition toward a coherent Responsible AI Use Policy anchored in three disclosure practices: explicit description of AI tool use in the introduction or methods section, in-text citation of AI tools with adequate specificity of contributor, model version, and date, and the submission of relevant AI interactions as supplementary appendices. For faculty and assessment design, the data support a parallel shift from product-based assessment of the final PDF thesis toward process-based assessment operationalised through mandatory viva voce defences, milestone logs, and continuous face-to-face supervisor checkpoints - an approach for which oral-viva

frameworks, project-based assessment models, and "Professional Conversation" safeguards are already articulated in the literature and translatable to the FTU-HCMC thesis-defence stage. Given the unreliability of AI-detection infrastructure and its disproportionate risk of false positives against multilingual writers, this shift is doubly warranted.

Four principal limitations bound the present claims: the data are self-reported and remain subject to social-desirability bias even after the indirect-questioning mitigation documented in Section 2.4; the sample is drawn from a single campus and cannot be extrapolated to broader Vietnamese or regional higher-education populations; the supporting tables do not disaggregate adoption by major, GPA, or gender; and the design is cross-sectional and therefore captures only a single temporal snapshot. Notwithstanding these boundaries, the study carries broader significance for Vietnamese and Southeast-Asian higher education: by surfacing the AI-integrity paradox as a structural feature of business-school GenAI adoption rather than a transient anomaly, the findings offer a reference point for institutional policy development across the regional cluster of applied-discipline programmes in which similar inversion patterns are emerging. Future research should pursue longitudinal designs of the kind piloted by Parker and colleagues and by Polyportis, multi-site comparative replication across Vietnamese institutions of varying prestige, and qualitative extension into supervisor perspectives that survey items cannot reliably detect.

References

- [1] Alhassany, H., & Faisal, F. (2018). Factors influencing the internet banking adoption decision in North Cyprus: An evidence from the partial least square approach of the structural equation modeling. *Financial Innovation*, 4(1). <https://doi.org/10.1186/s40854-018-0111-3>
- [2] Andrade, C. (2020). The inconvenient truth about convenience and purposive samples. *Indian Journal of Psychological Medicine*, 43(1), 86–88.
- [3] <https://doi.org/10.1177/0253717620977000>
- [4] Awdry, R., & Groves, A. (2023). Why they do and why they don't: A combined criminological approach to understanding assignment outsourcing in higher education. *International Journal for Educational Integrity*, 19(1). <https://doi.org/10.1007/s40979-023-00126-3>
- [5] Byrne, D. (2021). A worked example of Braun and Clarke's approach to reflexive thematic analysis. *Quality & Quantity*, 56(3), 1391–1412. <https://doi.org/10.1007/s11135-021-01182-y>
- [6] Cerri, J., Davis, E., Verissimo, D., & Glikman, J. A. (2021). Specialized questioning techniques and their use in

- conservation: A review of available tools, with a focus on methodological advances. *Biological Conservation*, 257, Article 109089.
- [7] <https://doi.org/10.1016/j.biocon.2021.109089>
- [8] Chan, C. K. Y., & Lee, K. K. W. (2023). The AI generation gap: Are Gen Z students more interested in adopting generative AI such as ChatGPT in teaching and learning than their Gen X and millennial generation teachers? *Smart Learning Environments*, 10(1).
- [9] <https://doi.org/10.1186/s40561-023-00269-3>
- [10] Chatto, M. J., & Logronio, F. (2025). Navigating AI in academia: Undergraduate experiences with ChatGPT and the redefinition of academic writing. *Journal of Education and Learning Reviews*, 2(4), 27–42. <https://doi.org/10.60027/jelr.2025.2061>
- [11] Dang, T. H. A., & Nguyen, V. T. T. (2025). Linking AI literacy, self-efficacy, attitudes, and achievement: A mixed-methods study on the moderating role of AI usage and study year. *Journal of Information Technology Education: Research*, 24, 45–45. <https://doi.org/10.28945/5689>
- [12] Danh, L. T. M., Nguyen, H. T., Tran, K. M. A., Dang, V. T., & Nguyen, B. K. H. (2024). Integrating TAM and UGT to explore students' motivation for using ChatGPT for learning in Vietnam. *Journal of Research in Innovative Teaching & Learning*, 19(1), 178–193. <https://doi.org/10.1108/jrit-05-2024-0116>
- [13] Dawadi, S., Shrestha, S., & Giri, R. A. (2021). Mixed-methods research: A discussion on its types, challenges, and criticisms. *Journal of Practical Studies in Education*, 2(2), 25–36. <https://doi.org/10.46809/jpse.v2i2.20>
- [14] Diloy, M. A. (2025). *The ChatGPT effect: Exploring generative AI use in student capstone projects and its implications for higher education* (pp. 486–490).
- [15] <https://doi.org/10.1109/waie67422.2025.11381028>
- [16] DiPietro, M. (2010). 14: Theoretical frameworks for academic dishonesty. *To Improve the Academy*, 28(1), 250–262. <https://doi.org/10.1002/j.2334-4822.2010.tb00606.x>
- [17] Drummond, M., & Dale, G. (2026). Generating a student-informed teaching and learning conceptual framework for GenAI in business schools: A case study. *Journal of University Teaching and Learning Practice*. <https://doi.org/10.53761/6jvtax83>
- [18] Durmaz, A., Dursun, İ., & Kabadayı, E. T. (2019). Mitigating the effects of social desirability bias in self-report surveys. In *Advances in library and information science (ALIS) book series* (pp. 146–185). SAGE Publishing. <https://doi.org/10.4018/978-1-7998-1025-4.ch007>
- [19] Gnambs, T., & Kaspar, K. (2014). Disclosure of sensitive behaviors across self-administered survey modes: A meta-analysis. *Behavior Research Methods*, 47(4), 1237–1259. <https://doi.org/10.3758/s13428-014-0533-4>
- [20] Ha, D. T. T. (2025). *The application of artificial intelligence in enhancing learning efficiency among Vietnamese university students*. Zenodo (CERN European Organization for Nuclear Research). <https://doi.org/10.5281/zenodo.17646113>
- [21] Hennink, M., & Kaiser, B. N. (2021). Sample sizes for saturation in qualitative research: A systematic review of empirical tests. *Social Science & Medicine*, 292, Article 114523. <https://doi.org/10.1016/j.socscimed.2021.114523>
- [22] Houk, A. A. H., Murphy, M., & Olsen, R. A. H. (2020). *Capstone courses and projects*. NC Digital Online Collection of Knowledge and Scholarship (The University of North Carolina at Greensboro).
- [23] Hung, N. M., & Goto, D. (2024). Unmasking academic cheating behavior in the artificial intelligence era: Evidence from Vietnamese undergraduates. *Education and Information Technologies*. <https://doi.org/10.1007/s10639-024-12495-4>
- [24] Hung, N. Q. (2025). Xu hướng lựa chọn đề tài luận văn tốt nghiệp của sinh viên chuyên ngành tiếng Trung Quốc. *Tạp chí Khoa học Trường Đại học Mở Hà Nội*, 66–66. <https://doi.org/10.59266/houjs.2024.513>
- [25] Ivankova, N. V., Creswell, J. W., & Stick, S. L. (2005). Using mixed-methods sequential explanatory design: From theory to practice. *Field Methods*, 18(1), 3–20. <https://doi.org/10.1177/1525822x05282260>
- [26] Ivanov, S., Soliman, M., Tuomi, A., Alkathiri, N. A., & Al-Alawi, A. N. (2024). Drivers of generative AI adoption in higher education through the lens of the Theory of Planned Behaviour. *Technology in Society*, 77, Article 102521. <https://doi.org/10.1016/j.techsoc.2024.102521>
- [27] Jin, Y., Yan, L., Echeverría, V., Gašević, D., & Martínez-Maldonado, R. (2024). Generative AI in higher education: A global perspective of institutional adoption policies and guidelines. *Computers and Education: Artificial Intelligence Trace*, 8, Article 100348. <https://doi.org/10.1016/j.caeai.2024.100348>
- [28] Malterud, K., Siersma, V., & Guassora, A. D. (2015). Sample size in qualitative interview studies. *Qualitative Health Research*, 26(13), 1753–1760.
- [29] <https://doi.org/10.1177/1049732315617444>
- [30] Naeem, M., Ozuem, W., Howell, K. E., & Ranfagni, S. (2023). A step-by-step process of thematic analysis to develop a conceptual model in qualitative research. *International Journal of Qualitative Methods*, 22. <https://doi.org/10.1177/16094069231205789>
- [31] Nakatani, K., & Jiang, Y. (2025). Understanding business students' GenAI usage and perception and associated implications. *Industry and Higher Education*.

- [32] <https://doi.org/10.1177/09504222251388167>
- [33] Nghia, T. L. H., Hoang, G., & Vo, P. Q. (2021). Business students' learning motivations and implications for sustainable implementations of at-home internationalized programs: Qualitative evidence from Vietnam. *Marketing Education Review*, 32(4), 329–341. <https://doi.org/10.1080/10528008.2021.2012701>
- [34] Nguyen, A. (2023). Unraveling EMI as a predictor of English proficiency in Vietnamese higher education: Exploring learners' backgrounds as a variable. *Pressto (Uniwersytetu Adama Mickiewicza)*, 13(2), 347–371. <https://doi.org/10.14746/ssllt.38278>
- [35] Nguyen, L. A. T., Nguyen, H., Tran, T. T. N., Le, M., & Zhang, Z. (2021). Factors influencing English proficiency: An empirical study of Vietnamese students. *Management Science Letters*, 1779–1786. <https://doi.org/10.5267/j.msl.2021.2.004>
- [36] Palos-Sánchez, P. R., Saura, J. R., Martín, M. Á. R., & Aguayo-Camacho, M. (2021). Toward a better understanding of the intention to use mHealth apps: Exploratory study. *JMIR mHealth and uHealth*, 9(9). <https://doi.org/10.2196/27021>
- [37] Passow, H. J., Mayhew, M. J., Finelli, C., Harding, T. S., & Carpenter, D. D. (2006). Factors influencing engineering students' decisions to cheat by type of assessment. *Research in Higher Education*, 47(6), 643–684. <https://doi.org/10.1007/s11162-006-9010-y>
- [38] Reginio, F. N., Marces, R. K. J. G., Lolong, H. M. G., & Niko, N. (2026). Redefining academic integrity in the age of AI. In *Advances in computational intelligence and robotics book series* (pp. 379–400). IGI Global. <https://doi.org/10.4018/979-8-3373-5601-3.ch016>
- [39] Ried, L., Eckerd, S., & Kaufmann, L. (2021). Social desirability bias in PSM surveys and behavioral experiments: Considerations for design development and data collection. *EUR Research Repository (Erasmus University Rotterdam)*, 28(1), Article 100743.
- [40] <https://doi.org/10.1016/j.pursup.2021.100743>
- [41] Son, N. T. K. (2023). Survey on graduation thesis writing of Vietnam's university students majoring in Chinese teaching as a foreign language. *Journal of Chinese Language and Culture Studies*, 2(1), 1–9. <https://doi.org/10.17977/um073v2i12023p1-9>
- [42] Stratton, S. J. (2021). Population research: Convenience sampling strategies. *Prehospital and Disaster Medicine*, 36(4), 373–374. <https://doi.org/10.1017/s1049023x21000649>
- [43] Taherdoost, H. (2016). Validity and reliability of the research instrument; How to test the validation of a questionnaire/survey in a research. *SSRN Electronic Journal*.
- [44] <https://doi.org/10.2139/ssrn.3205040>
- [45] Taslim, M., Putra, R. P., Daulay, N., & Bulut, S. (2025). Academic cheating with generative AI in higher education: An extended model of the theory of planned behavior with motivational antecedents. *Jurnal Psikologi*, 52(3), 313–313. <https://doi.org/10.22146/jpsi.107932>
- [46] Thanapongporn, A., Saengchote, K., & Gowanit, C. (2023). Factors shaping Thai millennials' low-carbon behavior: Insights from extended theory of planned behavior. *HighTech and Innovation Journal*, 4(3), 482–504. <https://doi.org/10.28991/hij-2023-04-03-02>
- [47] Thao, V. T. T. (2025). The impact of artificial intelligence on graduation thesis writing: A case study of English majors at a Vietnamese public university. *International Journal on Studies in English Language and Literature*, 13(9), 43–53. <https://doi.org/10.20431/2347-3134.1309005>
- [48] Tien, D. T. A. N. H., & Ngoc, L. D. M. D. N. M. (2021). Sustainable development of higher education: A case of business universities in Vietnam. *Journal of Hunan University*, 47(12). <https://johuns.net/index.php/journal/article/view/488/485.html>
- [49] Toyon, M. A. S. (2021). Explanatory sequential design of mixed methods research: Phases and challenges. *International Journal of Research in Business and Social Science*, 10(5), 253–260. <https://doi.org/10.20525/ijrbs.v10i5.1262>
- [50] Tudy, R. A., Chin, C. K., Gabales, B. J., & Tudy, I. (2025). Welcoming or banning ChatGPT in higher education: Insights from Indonesia, Malaysia, and the Philippines. *The Palawan Scientist*, 17(2), 13–22. <https://doi.org/10.69721/tps.j.2025.17.2.02>
- [51] Wang, H., Dang, A. K., Wu, Z., & Mac, S. (2024). Generative AI in higher education: Seeing ChatGPT through universities' policies, resources, and guidelines. *Computers and Education: Artificial Intelligence*, 7Trace, Article 100326.
- [52] <https://doi.org/10.1016/j.caeai.2024.100326>
- [53] Westland, H., Vervoort, S. C. J. M., Kars, M. C., & Jaarsma, T. (2024). Interviewing people on sensitive topics: Challenges and strategies. *KTH Publication Database DiVA (KTH Royal Institute of Technology)*, 24(3), 488–493. <https://doi.org/10.1093/eurjcn/zvae128>
- [54] Xiao, P., Chen, Y., & Bao, W. (2023). Waiting, banning, and embracing: An empirical analysis of adapting policies for generative AI in higher education. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.4458269>
- [55] Zhang, X., Hu, X., Sun, Y. Z., Li, L., Deng, S., & Xiaowen, C. (2025). Integrating AI literacy with the TPB-TAM framework to explore Chinese university students' adoption of generative AI. *Behavioral Sciences*, 15(10), Article 1398. <https://doi.org/10.3390/bs15101398>