

Original Research

Drivers of Ethical LLM Usage in Chinese Higher Education: Integrating Social Cognitive and Rational Choice Theories

Ci Zhang¹, Wilson Cheong Hin Hong², Na Cai³, Yunfeng Zhang^{4,*}, Xiaoshu Xu^{1,*}

¹School of Foreign Studies, Wenzhou University, 325000 Wenzhou, Zhejiang, China

²Faculty of Innovation Hospitality Management/Faculty of Creative Tourism and Intelligent Technologies, Macao University of Tourism, 999078 Macao, China

³School of Education Science, Guangdong Polytechnic Normal University, 510665 Guangzhou, Guangdong, China

⁴Faculty of Languages and Translation, Macao Polytechnic University, 999078 Macao, China

*Correspondence: zhangyunfeng@mpu.edu.mo (Yunfeng Zhang); Lisaxu@wzu.edu.cn (Xiaoshu Xu)

Submitted: 21 March 2026 Revised: 25 July 2026 Accepted: 30 July 2026 Published: 28 August 2026

Abstract

The rapid integration of large language models (LLMs) in higher education presents significant opportunities for personalized learning alongside profound ethical challenges, such as academic dishonesty and cognitive over-reliance. Moving beyond technology acceptance models, this study investigates the structural drivers of ethical LLM usage among Chinese university students by integrating social cognitive theory (SCT) and rational choice theory (RCT). Utilizing an adapted, culturally validated multidimensional scale, survey data from 519 students were analyzed using confirmatory factor analysis (CFA) and partial least squares structural equation modeling (PLS-SEM). The findings reveal a stark contrast between institutional and interpersonal influences: while abstract academic ethical climates and mere policy awareness fail to ensure ethical behavior, direct ethical supervision shows a strong positive association with students' moral consciousness and ethical use (SCT). Conversely, prolonged use of LLMs is negatively associated with ethical behavior, suggesting that as students gain technical proficiency, their rational cost-benefit calculations often tip toward unethical shortcuts unless credible deterrents—such as detection probability and sanction severity—are present (RCT). This study contributes to the theoretical understanding of AI ethics by demonstrating that ethical technology use is a negotiated outcome between social modeling and rational risk assessment. Practically, it urges institutions to transition from passive policy publishing to active supervisor training, balanced oversight, and the cultivation of peer-level moral self-efficacy.

Keywords: ethical practices of large language models (LLMs); Chinese higher education; social cognitive theory; rational choice theory; structural equation modeling

Introduction

Large language models (LLMs), powered by advanced natural language processing (NLP) and extensive pre-trained data, have revolutionized education by offering human-like understanding and performance in various tasks (Kasneci et al., 2023). Tools like ChatGPT and BERT enhance learning efficacy through personalized experiences (Farrokhnia et al., 2024; Hong, 2026) and adaptive assessments, while Chinese-developed LLMs such as Wenxin Yiyan and Deepseek address specific needs like mathematical and logical problem-solving (Neha & Bhati, 2025; Tian et al., 2024). These tools are increasingly integrated into higher education to improve teaching and learning outcomes, particularly in China, where the demand for comprehensive practical talents is high and AI adoption rates among students are rapidly accelerating (Song, 2024; Sun et al., 2025).

Despite their benefits, LLMs pose significant ethical challenges. For instance, their reliance on biased training data can mislead teachers and learners, while over-reliance on these tools may lead to unethical practices, such as academic dishonesty, plagiarism, and the generation of fabricated information or hallucinations (Dowling & Lucey, 2023; Kasneci et al.,



2023; Yan et al., 2025). Additionally, the absence of comprehensive and standardized ethical frameworks for LLM use in education intensifies these challenges by placing students and educators in the difficult position of integrating AI technologies without adequate guidance, while simultaneously avoiding the development of a counterproductive culture characterized by excessive monitoring and policing. (Gorelick & McDonald, 2023; Hong, 2026). While some Chinese universities, such as Chinese University of Hong Kong (CUHK) and Fudan, have tried addressing these concerns by establishing AI usage guidelines, a comprehensive understanding of the underlying cognitive and environmental factors influencing ethical LLM use remains underexplored (CUHK, 2023; UCLA, 2023).

As LLM usage increasingly permeates higher education, ethical LLM integration is also relevant to Whole Schooling and inclusive educational practice. Ensuring equitable access to learning technologies is important because students differ in linguistic background, academic preparedness, digital literacy, and access to learning support. In this sense, ethical LLM use should not be understood only as a matter of academic integrity, but also as part of responsible student support. When guided appropriately, LLMs can provide linguistic scaffolding, cognitive support, and accessibility-related assistance for students who may otherwise face barriers to participation. At the same time, ethical integration requires educators to design authentic assessments that value students' learning processes, transparent attribution, and responsible human–AI collaboration. Therefore, institutional policies should not simply police students' AI use; rather, they should guide students, teachers, and institutions toward a caring academic community in which technology supports fairness, trust, and holistic student development (Kong & Zhu, 2025).

Existing research has primarily focused on technology acceptance and adoption intentions, with limited attention to the structural drivers of ethical behavior post-adoption in higher education, particularly in China (Rout et al., 2024; Shahzad et al., 2024). Furthermore, studies often lack robust, integrated theoretical frameworks to analyze ethical practices. This study addresses these gaps by integrating social cognitive theory (SCT) and rational choice theory (RCT) to examine how environmental modeling and rational risk-benefit calculations jointly shape Chinese students' cognition and behavior regarding LLM usage. Using an adapted version of Mvondo et al. (2024)'s ethical use scale, this research aims to:

1. Adapt a multidimensional ethical use scale tailored to the Chinese higher education context from perspectives of ethical climate and supervision; awareness of AI chatbot usage policies; moral consciousness and ethical chatbot self-efficacy; AI writing detection probability and sanction severity.
2. Examine the structural relationships of different dimensions corresponding to environmental, technological, individual, and risk factors and their association with Chinese students' ethical use of LLMs.
3. Provide evidence-based, actionable recommendations for institutions, policymakers, teachers, and students accordingly.

By addressing these objectives, this study seeks to advance the theoretical understanding of AI ethics and inform the development of ethical frameworks and institutional policies governing the use of LLMs in Chinese higher education, with the aim of balancing the pedagogical benefits of AI technologies with the imperative to uphold academic integrity and promote meaningful learning.

Literature Review

LLM Applications in Higher Education

LLMs are a subset of generative AI based on advanced neural network architectures that enable interactions from input sequences' attention scores to create subsequent predicting content (Lv, 2023). This capability is becoming increasingly relevant in higher education, where approximately 86% of education organizations reported using generative AI by 2025 (Microsoft Education, 2025). Studies indicate that nearly a third of college students have adopted these tools to complete school tasks (Song, 2024). Using varied frameworks (e.g., Technology Acceptance Model, Unified Theory of Acceptance and Use of Technology, Expectation Confirmation Model), studies have highlighted the flexibility, accessibility, and interactivity of AI chatbots, which make them valuable tools for personalized learning and adaptive assessments (Schei et al., 2024; Tian et al., 2024).

Furthermore, studies emphasize the potential of LLMs, such as ChatGPT, to serve as powerful tools for language

learning (Hong, 2023), writing support (Barrot, 2023; Dergaa et al., 2023), and even exam preparation, where early stages of ChatGPT already demonstrated human-comparable performances (Kung et al., 2023). In the context of Chinese higher education, LLMs serve as critical assistive tools for English as a Foreign Language (EFL) thesis writing, providing alternative sources of language support for mechanics and lexical choices (Hong & Litwin, 2025). Disciplinary differences also shape adoption patterns; for instance, engineering students exhibit the highest daily usage rates for computational problem-solving (Sun et al., 2025). Moreover, pre-service teachers in China often conceptualize AI through the metaphor of a “seed”, viewing it as a supportive, developmental technology with the potential for massive future growth (Hu et al., 2025). Nevertheless, LLMs’ potential to assist human learners across a wide range of areas comes at the cost of new challenges and ethical concerns.

Challenges and Ethical Concerns

Despite these benefits, LLMs pose significant challenges to academic integrity and learning efficacy. While they can assist students in writing, feedback, and information searching, their inappropriate or excessive use may reduce students’ opportunities to engage independently in planning, reasoning, drafting, and revision (Susnjak, 2022). In higher education, such concerns are particularly relevant because students may use LLMs not only for language support or idea generation, but also for unauthorized text production, fabricated references, data fabrication, or assessment-related shortcuts. This form of over-reliance may encourage students to copy, adapt, or submit AI-generated outputs without fully understanding the underlying content or acknowledging the role of AI (Yan et al., 2025). Specific to Chinese higher education, research on tools such as Open-Buddy has noted that while AI can help students produce more fluent sentences, it may also generate incoherent logic, fabricated information, non-existent references, and hallucinated content (Hong & Litwin, 2025; Li et al., 2023; Yan et al., 2025). Therefore, the key issue is not simply whether students use LLMs, but how they make ethical decisions when using them in academic tasks.

Meanwhile, the use of LLMs raises important ethical questions about accountability, data privacy, algorithmic bias, and transparency. The widespread use of these tools has blurred traditional boundaries of originality, resulting in an increase in plagiarism and necessitating a paradigm shift in how academic dishonesty is defined in the digital age (Hutson, 2024; Song, 2024). While tools like ChatGPT can convincingly mimic human interaction, they lack a genuine understanding of ethics and attribution (Bass, 2022), making it easy for students to cheat the system by feeding assessment information into the AI and tweaking the output (Dwivedi et al., 2023).

While institutions have attempted to implement detection software, these tools are notoriously unreliable and frequently exhibit bias against non-native English speakers (Hong, 2026). The lack of formal guidelines at many institutions—with less than 10% having developed formal guidance as of 2023—leads to confusion and a counter-productive “policing culture” that can damage the relational trust between students and faculty (Song, 2024).

Factors Influencing Ethical LLM Usage and the Research Gap

Research has identified several factors influencing the use of LLMs. General adoption is often driven by task types, perceived social norms (Zhang et al., 2025), and cognitive/affective factors such as perceived usefulness and functional value (Rout et al., 2024). In China, attitudes are significantly shaped by direct classroom exposure (Hu et al., 2025) and disciplinary demands (Sun et al., 2025).

A key gap in the current literature is that many studies focus on technology acceptance and adoption intentions, while fewer studies examine the structural factors associated with ethical behavior after adoption. Although problems such as plagiarism, fabricated references, hallucinated content, and unauthorized AI-generated writing have been widely discussed, the cognitive, environmental, technological, and risk-related factors associated with students’ ethical decisions remain insufficiently understood. Therefore, promoting ethical LLM use requires moving beyond general adoption models and punitive rules to examine how students’ ethical behavior is shaped by supervisory guidance, institutional norms, moral consciousness, ethical self-efficacy, detection probability, and sanction severity (Mvondo et al., 2024).

Recent empirical studies have begun to address this by proposing multidimensional models grounded in SCT and RCT. For example, Mvondo et al. (2024) suggest that ethical supervision, moral consciousness, ethical chatbot self-efficacy, and the perceived probability of detection collectively determine usage outcomes. The study mainly investigated four dimensions:

- **Ethical Climate and Supervision (Environmental Factors):** Ethical climate represents shared norms and values within an organization that shape ethical decision-making (Martin & Cullen, 2006). In higher education, it reflects expectations guiding students' LLM usage, while ethical supervision, through close student-supervisor relationships, fosters ethical behavior (Nejati & Shafaei, 2018).
- **Awareness of AI Chatbot Usage Policies (Technological Factors):** Awareness of LLM usage policies is critical for ensuring ethical behavior. This dimension examines students' understanding of institutional guidelines and their role in shaping responsible LLM usage.
- **Moral Consciousness and Ethical Chatbot Self-Efficacy (Individual Factors):** Moral consciousness, influenced by social norms, experiences, and culture, guides ethical decision-making. Ethical chatbot self-efficacy refers to confidence in making ethical choices when using LLMs, both of which shape responsible usage.
- **AI Writing Detection Probability and Sanction Severity (Risk Factors):** Perceptions of detection tools' effectiveness and the severity of penalties for misconduct influence students' willingness to engage in unethical practices, with higher risks and penalties deterring misuse.

While this ethical use scale provides a comprehensive framework for understanding the factors influencing ethical LLM usage (see Figure 1 for Mvondo et al. (2024) model), a significant gap remains in applying and validating this framework within non-Western educational contexts. Chinese higher education presents a unique and impactful ecosystem characterized by high power distance (where supervisor influence is paramount), collectivist learning cultures (Hofstede, 2011), and a rapid, large-scale rollout of domestic and international AI tools amidst lagging institutional policy formalization (Hong, 2026). Furthermore, while commentary articles are rich (e.g., Gallent-Torres et al., 2023; Hong, 2026; Hutson, 2024;), existing empirical research provides limited insight into how rational risk assessment and social modeling interact to influence students' ethical use of LLMs. Therefore, by adapting Mvondo et al. (2024) dimensions to the Chinese higher education context, this research provides a much-needed empirical investigation into how environmental, technological, individual, and risk factors structurally interact to shape the ethical usage of LLMs, ultimately offering evidence-based guidelines for institutions navigating the AI transition.

Theoretical Framework

Integration of SCT and RCT

To comprehensively understand the ethical usage of LLMs, this study integrates SCT and RCT. Rather than operating in isolation, these two theories provide a complementary lens through which to view student behavior in complex, technology-rich environments. Compared with frameworks that focus mainly on behavioral intention or individual self-regulation, SCT is appropriate for the present cross-sectional survey because it provides a relational lens for examining how personal factors, environmental conditions, and ethical behavior are structurally associated at a given point in time. Its concept of reciprocal determinism helps explain why ethical LLM use should be examined not only as an individual choice, but also as a behavior embedded in supervisory guidance, institutional norms, moral consciousness, and ethical self-efficacy. Therefore, SCT was adopted together with RCT to capture both socially embedded ethical learning and students' risk–benefit calculations in LLM use.

SCT, developed by Bandura (1986), posits a triadic reciprocal relationship among personal factors (e.g., moral consciousness), environmental factors (e.g., ethical climate and supervision), and behavior. In the context of LLM usage, SCT suggests that a student's internal ethical compass and self-efficacy are continuously shaped by the academic environment and the modeling of behaviors by supervisors and peers (Mvondo et al., 2024).

Conversely, RCT focuses on the immediate decision-making process, positing that individuals act to maximize benefits while minimizing risks (Becker, 1968). When deciding whether to use an LLM unethically (e.g., for plagiarism), students weigh the perceived benefits (e.g., saving time, getting a better grade) against the perceived risks (e.g., AI detection probability, sanction severity).

This study argues that SCT and RCT interact dynamically: the environmental and personal factors outlined in SCT serve as the foundational lens through which students conduct their RCT-based cost-benefit analyses. For instance, a strong ethical climate and supportive supervision (SCT) can elevate a student's moral consciousness, which in turn increases the subjective cost of unethical behavior in their rational calculations (RCT). By combining these frameworks, we can better

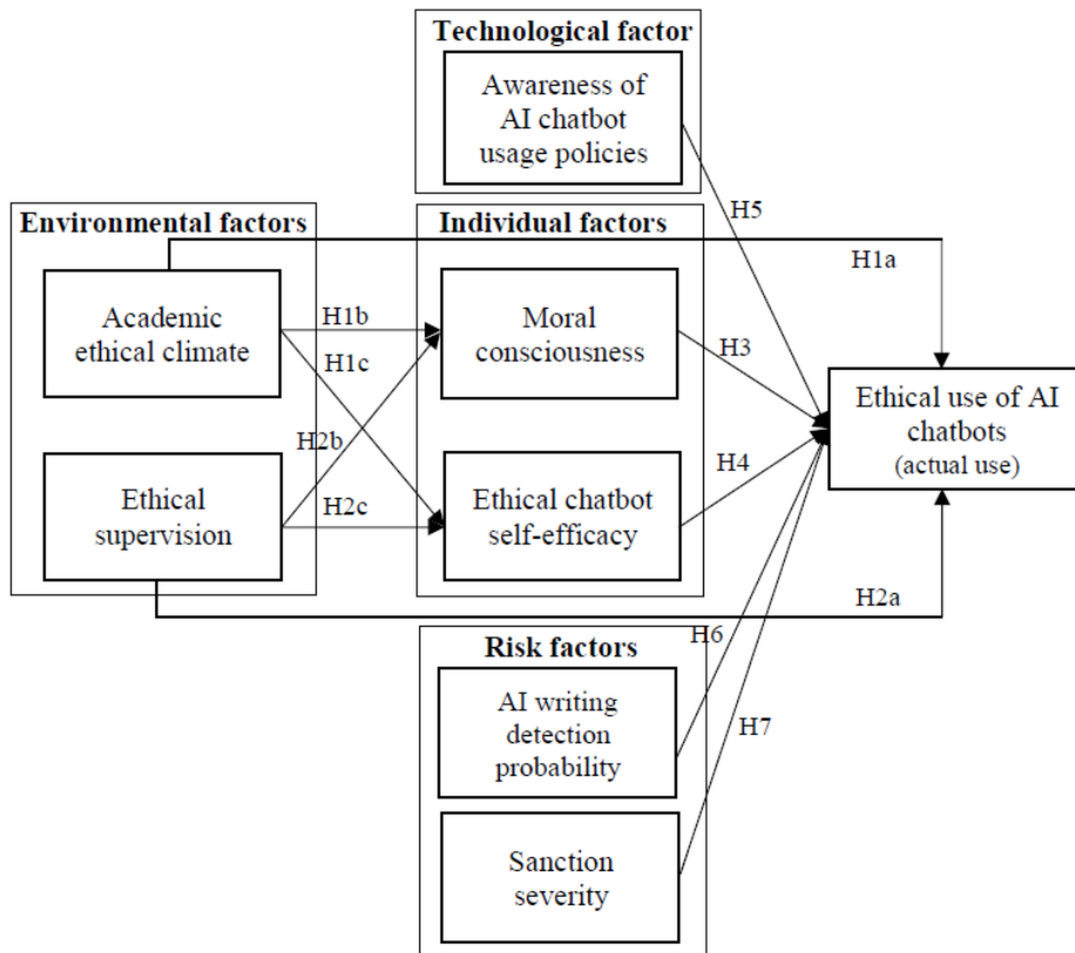


Figure 1. Mvondo et al.'s (2024) model. Adapted from "Exploring the ethical use of LLM chatbots in higher education," by G. F. N. Mvondo, B. Niu, and S. Eivazinezhad, 2024, SSRN Electronic Journal, p.17. <https://doi.org/10.2139/ssrn.4548263>.

map the structural relationships between environment, individual cognition, risk assessment, and ultimate ethical behavior.

Research Questions and Hypotheses

This study addresses the following research question: Which environmental, technological, individual, and risk factors are structurally associated with students' ethical usage of LLMs in Chinese higher education?

Drawing on the integrated SCT-RCT framework and adapting Mvondo et al. (2024) model (see Figure 1), we propose the following hypotheses. Note that in alignment with our cross-sectional design, these hypotheses posit structural associations rather than strict causal determinations.

- Environmental Factors (SCT):

- H1a: Academic ethical climate is positively associated with ethical use.
- H1b: Academic ethical climate is positively associated with moral consciousness.
- H1c: Academic ethical climate is positively associated with ethical chatbot self-efficacy.
- H2a: Ethical supervision is positively associated with ethical use.
- H2b: Ethical supervision is positively associated with moral consciousness.
- H2c: Ethical supervision is positively associated with ethical chatbot self-efficacy.

- Individual Factors (SCT & RCT intersection):
 - H3: Moral consciousness is positively associated with ethical use.
 - H4: Ethical chatbot self-efficacy is positively associated with ethical use.
- Technological Factors (RCT):
 - H5: Awareness of chatbot usage policies is positively associated with ethical use.
- Risk Factors (RCT):
 - H6: AI writing detection probability is positively associated with ethical use.
 - H7: Sanction severity is positively associated with ethical use.

Methodology

Sample and Sampling

Graduates and undergraduates from 7 universities in a Southern coastal city in China from 20 September to 27 October 2025, representing diverse disciplines ranging from computing and engineering to liberal arts, participated in the survey. Our study utilized convenience sampling, supplemented by snowball sampling to ensure a sufficient sample size for structural equation modeling (Seidman, 2019). The targeted participants were required to have a minimum of one year of experience utilizing online learning technology. A total of 550 individuals consented via the Wenjuanxing platform and completed the survey. The research protocol was approved by the first author's university academic committee (IRB#20260226), and all personal identifiers were replaced with unique codes to protect participant confidentiality.

Research Instrument and Validation

This study employed a quantitative approach using a 38-item questionnaire adapted from Mvondo et al. (2024). These five-point scale questions measured Ethical Climate, Ethical Supervision, Moral Consciousness, Ethical Chatbot Self-Efficacy, Awareness of Chatbot Usage Policies, AI Writing Detection Probability, Sanction Severity, and the dependent variable, Ethical Use of AI chatbots, measured using a single-item (referencing Mvondo et al. (2024)). This was to reduce defensive case-based self-reporting. However, the question was modified to define specific unethical AI practices (e.g., ghost-writing, data & sources fabrication, cheating, and unauthorized uploads) to clearly anchor students' understanding.

To ensure the cultural and contextual appropriateness of the survey instrument, Mvondo et al. (2024) original English scale underwent a rigorous adaptation process. First, a forward-backward translation method was employed by two bilingual researchers to translate the items into Mandarin Chinese and back into English, ensuring semantic equivalence. Second, an expert panel consisting of three experienced IT researchers reviewed the translated items for content validity, adjusting phrasing to align with Chinese academic terminology. Finally, a pilot test was conducted with 30 undergraduate students to verify the clarity and readability of the questions (see Appendix Table 8) before the full-scale deployment.

Data Handling and Analytic Strategy

Upon initial data cleaning procedures, abnormal responses, including straight-lining or overly short completion times, were removed, resulting in 519 valid responses (see Table 1 for the demographic distribution). To ensure a coherent psychometric validation strategy, data analysis was conducted in three phases using SPSS 29.0 (IBM Corp., NY, USA) and SmartPLS4.1 (SmartPLS GmbH, Monheim am Rhein, Germany). As the Chinese context is novel, though not a research objective here, an initial Exploratory Factor Analysis was performed (see Appendix for results), alongside regressions to test for the effects of the demographic factors. Through these preliminary analyses, "length of chatbot use" emerged as a significant predictor of ethical behavior. Consequently, this variable—coded ordinally into five categories (see Table 1), was incorporated into the structural model as a control variable. This allowed us to account for students' potential proficiency levels while ensuring that its inclusion did not alter the interpretation of our primary hypothesized paths.

We utilized confirmatory factor analysis (CFA) to rigorously validate the measurement model, assessing construct reliability (CR) and average variance extracted (AVE) for the adapted Chinese scale. Finally, Structural Equation Modeling (SEM) was employed to test the hypothesized structural associations between the latent variables, providing a robust

evaluation of the integrated SCT-RCT framework.

Table 1. Demographic information of the respondents.

Variable	Categories	n	%
Age	<20	323	62.2%
	20–25	166	32.0%
	26–30	13	2.5%
	>30	17	3.3%
Gender	Male	175	33.7%
	Female	344	66.3%
Highest education	Secondary education	11	2.1%
	Diploma/Associate degree	57	11.0%
	Bachelor degree	397	76.5%
	Master's degree	36	6.9%
	Doctoral degree	11	2.1%
	Other	7	1.3%
Discipline	Social science	288	55.5%
	Natural science	33	6.3%
	Technology & engineering	167	32.2%
	Arts & humanities	31	6.0%
Length of chatbot usage	<3 months	157	30.3%
	3–6 month	30	5.8%
	6–9 months	35	6.7%
	9–12 months	50	9.6%
	>12 months	247	47.6%

Results

Confirmatory Factor Analysis

The measurement model was assessed using partial least squares structural equation modeling (PLS-SEM) in accordance with established guidelines for reflective constructs (Hair et al., 2019; Hair & Sarstedt, 2022). First, CFA was performed. All latent variables demonstrated strong psychometric properties, supporting convergent validity, internal consistency reliability, and discriminant validity.

Convergent validity and reliability were established through examination of indicator loadings (all > 0.70; see Table 2), average variance extracted (AVE), composite reliability (ρ_c), and Cronbach's alpha. As presented in Table 2, all constructs exceeded recommended thresholds: Cronbach's alpha values ranged from 0.823 to 0.948, composite reliability (ρ_c) from 0.919 to 0.959, and AVE from 0.708 to 0.901 (all > 0.50; Hair et al., 2019). These results indicate that the items adequately converged on their respective constructs and that the constructs possessed sufficient internal consistency and convergent validity.

Table 2. Construct reliability and convergent validity.

Construct	Items	Cronbach's α	CR (ρ_a)	CR (ρ_c)	AVE
AI writing detection probability	2	0.823	0.823	0.919	0.85
Awareness of chatbot usage policies	3	0.909	0.913	0.943	0.847
Ethical chatbot self-efficacy	5	0.947	0.947	0.959	0.826
Ethical climate	5	0.936	0.949	0.952	0.798
Ethical supervision	10	0.948	0.949	0.956	0.708
Moral consciousness	6	0.932	0.933	0.946	0.746
Sanction severity	2	0.891	0.894	0.948	0.901

AVE, average variance extracted; CR, composite reliability.

Discriminant validity was assessed using the Heterotrait-Monotrait ratio (HTMT) criterion and the Fornell-Larcker criterion. All HTMT values were below the conservative threshold of 0.85 (Hair et al., 2019) (see Table 3), with the highest inter-construct HTMT being 0.797 (Ethical Supervision ↔ Moral Consciousness). The Fornell-Larcker criterion was also satisfied: the square root of each construct's AVE (displayed on the diagonal in Table 4) exceeded all corresponding off-diagonal correlations. These results confirm that the constructs were empirically distinct.

Table 3. HTMT criterion for discriminant validity.

Construct	1	2	3	4	5	6	7	8
1. AI writing detection probability								
2. Awareness of chatbot usage policies	0.321							
3. Ethical chatbot self-efficacy	0.303	0.783						
4. Ethical climate	0.345	0.575	0.525					
5. Ethical supervision	0.349	0.657	0.616	0.774				
6. Ethical use	0.297	0.362	0.430	0.329	0.457			
7. Length of chatbot use	0.153	0.187	0.087	0.230	0.155	0.230		
8. Moral consciousness	0.400	0.689	0.683	0.615	0.797	0.486	0.116	
9. Sanction severity	0.502	0.543	0.505	0.465	0.523	0.383	0.117	0.586

HTMT, Heterotrait-Monotrait ratio.

Table 4. Fornell–Larcker criterion for discriminant validity.

Construct	1	2	3	4	5	6	7	8	9
1. AI writing detection probability	0.922								
2. Awareness of chatbot usage policies	0.279	0.92							
3. Ethical chatbot self-efficacy	0.268	0.725	0.909						
4. Ethical climate	0.304	0.532	0.499	0.893					
5. Ethical supervision	0.309	0.613	0.585	0.736	0.841				
6. Ethical use	0.269	0.346	0.418	0.322	0.446	1			
7. Length of chatbot use	0.139	0.178	0.084	0.216	0.151	−0.23	1		
8. Moral consciousness	0.35	0.635	0.641	0.582	0.75	0.47	0.112	0.864	
9. Sanction severity	0.428	0.489	0.464	0.428	0.483	0.362	0.112	0.534	0.949

Note. Square roots of AVE are on the diagonal (bold). Off-diagonal values are inter-construct correlations. AVE, average variance extracted.

Additionally, cross-loadings (not tabulated for brevity but examined) confirmed that each indicator loaded most strongly on its intended construct (loadings >0.70) and weakly on others (differences >0.10 in most cases), further supporting discriminant validity.

Model fit indices for the saturated and estimated models were acceptable in the PLS-SEM context (Henseler & Ray, 2016). The standardized root mean square residual (SRMR) was 0.044 for the saturated model and 0.084 for the estimated model (both <0.10, with the saturated model indicating good fit). Other discrepancy measures (d_{ULS} , d_G) and fit indices ($NFI \approx 0.86$ saturated) were consistent with adequate approximate model fit for PLS-SEM applications.

Structural Modeling

Hypothesized Effects

Then, using the same PLS-SEM with 5000 bootstrap resamples (Hair & Sarstedt, 2022). All reported path coefficients are standardized. The model explained 56.2% of the variance in moral consciousness ($R^2 = .562$), 35% of the variance in ethical chatbot self-efficacy ($R^2 = .35$), and 34.4% of the variance in ethical use of AI chatbots ($R^2 = .344$). These values indicate moderate to substantial explanatory power. Results of the hypothesis tests are presented in Table 5.

Table 5. Direct effect results.

H	Path	β	t	p	2.5% CI	97.5% CI	f^2	Decision
Environmental factors (SCT)								
H1a	Ethical climate → Ethical use	-0.01	0.08	0.937	-0.099	0.111	0	Not supported
H1b	Ethical climate → Moral consciousness	0.066	1.143	0.253	-0.051	0.176	0.01	Not supported
H1c	Ethical climate → Ethical chatbot self-efficacy	0.149	2.311	0.021	0.02	0.271	0.02	Supported
H2a	Ethical supervision → Ethical use	0.21	2.85	0.004	0.087	0.558	0.02	Supported
H2b	Ethical supervision → Moral consciousness	0.701	14.296	<.001	0.603	0.795	0.52	Supported
H2c	Ethical supervision → Ethical chatbot self-efficacy	0.475	8.05	<.001	0.357	0.586	0.16	Supported
Individual factors								
H3	Moral consciousness → Ethical use	0.292	5.449	<.001	0.189	0.397	0.02	Supported
H4	Ethical chatbot self-efficacy → Ethical use	0.172	3.305	0.001	0.071	0.273	0.02	Supported
Technological factor (RCT)								
H5	Awareness of chatbot usage policies → Ethical use	0.005	0.094	0.925	-0.102	0.107	0	Not supported
—	Length of chatbot use → Ethical use	-0.306	9.395	<.001	-0.37	-0.241	0.15	Significant (negative)
Risk factors (RCT)								
H6	AI writing detection probability → Ethical use	0.32	2.67	0.008	0.089	0.56	0.02	Supported
H7	Sanction severity → Ethical use	0.109	2.261	0.024	0.015	0.204	0.01	Supported

Note. Effect sizes (f^2): 0.02 = small, 0.15 = medium, 0.35 = large [Cohen \(1988\)](#). SCT, social cognitive theory; RCT, rational choice theory.

Indirect Effects

The total indirect effect of ethical supervision on ethical use was significant ($\beta = 0.205$, $t = 4.97$, $p < .001$), whereas the total indirect effect of ethical climate was marginally significant ($\beta = 0.026$, $t = 1.7$, $p = .074$) referencing 95% CI (Hair & Sarstedt, 2022). See Table 6 for detailed results.

Table 6. Indirect effect results.

Mediation path	β	t	p	2.5% CI	97.5% CI	Decision
Ethical climate → Moral consciousness → Ethical use	0.019	1.101	0.271	−0.014	0.056	No Mediation
Ethical supervision → Moral consciousness → Ethical use	0.205	4.969	<.001	0.126	0.289	Mediation
Ethical climate → Ethical chatbot self-efficacy → Ethical use	0.026	1.789	0.074	0.003	0.057	Mediation*
Ethical supervision → Ethical chatbot self-efficacy → Ethical use	0.082	3.067	0.002	0.031	0.135	Mediation

*based on the 95% bootstrap confidence interval excluding zero.

Model Evaluation

The structural model demonstrated acceptable predictive power. Collinearity was not a concern (inner VIF values < 5, assumed from standard PLS-SEM diagnostics) (see Table 7). The model fit indices from the measurement model (SRMR saturated = 0.044; estimated = 0.084) remained acceptable when the structural paths were included.

Table 7. Inner VIF.

Path	VIF
AI writing detection probability → Ethical use	1.27
Awareness of chatbot usage policies → Ethical use	2.47
Ethical chatbot self-efficacy → Ethical use	2.41
Ethical climate → Ethical chatbot self-efficacy	2.18
Ethical climate → Moral consciousness	2.18
Ethical supervision → Ethical chatbot self-efficacy	2.18
Ethical supervision → Moral consciousness	2.18
Length of chatbot use → Ethical use	1.05
Moral consciousness → Ethical use	2.11
Sanction severity → Ethical use	1.63

VIF, variance inflation factor.

In summary, eight of the eleven hypothesized direct relationships (H1c, H2a, H2b, H2c, H3, H4, H6, H7) were supported. H1a, H1b, and H5 were not supported. Ethical supervision emerged as the strongest predictor overall, exerting both direct and substantial indirect effects through moral consciousness and ethical chatbot self-efficacy. The negative relationship between length of chatbot use and ethical use was noteworthy and warrants further investigation in future studies.

These findings partially replicate the original model (Mvondo et al., 2024) while revealing meaningful contextual differences, particularly regarding the limited role of policy awareness and the emergence of a direct effect from ethical supervision to ethical use. The final updated model is illustrated in Figure 2.

Discussion

This study integrated SCT and RCT to examine factors associated with ethical LLM usage among Chinese university students. The findings partially support previous models of AI ethics (Mvondo et al., 2024), while also showing several context-specific patterns. Specifically, direct ethical supervision and perceived external risks were more clearly associated with ethical use than broad institutional ethical climate or policy awareness.

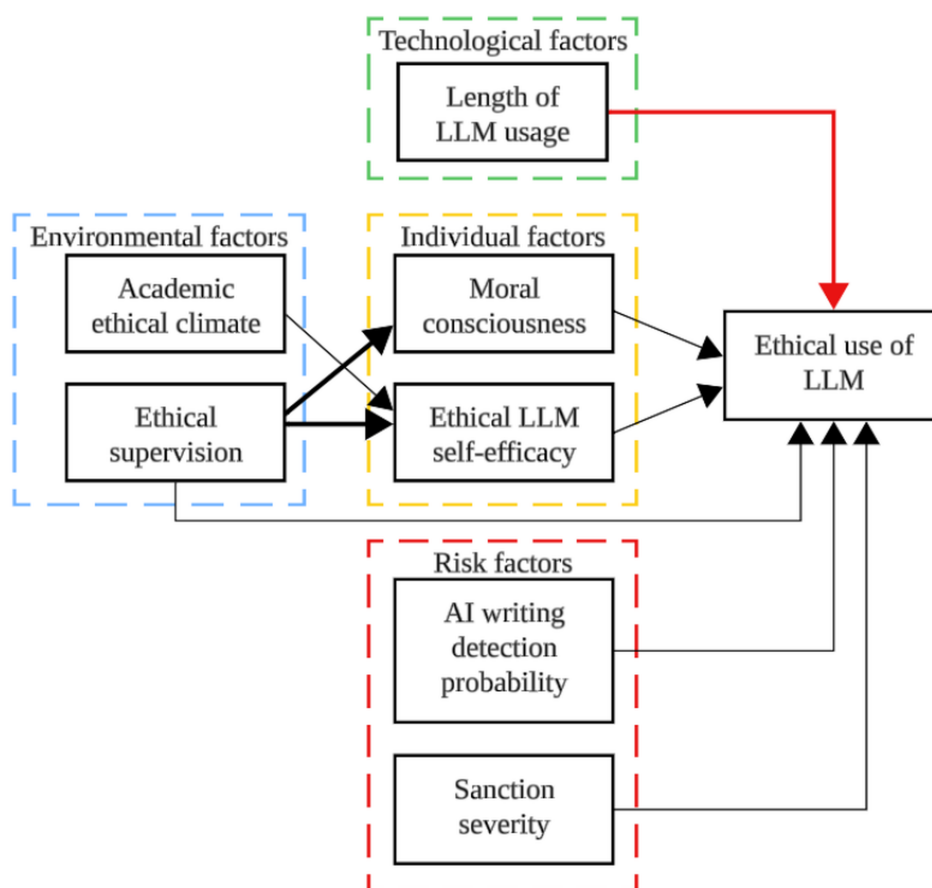


Figure 2. Revised model based on current results. → moderate/strong positive effect, → positive effect, → moderate negative effect. LLM, large language model.

The Primacy of Ethical Supervision Over Institutional Climate

A striking finding of this study is the stark contrast between the impact of the academic ethical climate and direct ethical supervision. While the broader institutional climate failed to significantly influence moral consciousness or ethical use, ethical supervision emerged as a robust predictor, displaying strong direct and indirect associations to students' ethical behavior.

This divergence aligns with the systemic characteristics of Chinese higher education, including relatively strong supervisor–student relationships and collectivist learning cultures (Hofstede, 2011). In this context, mentors and supervisors may play an important role in guiding students' academic practices and ethical decision-making (Li et al., 2025). This is consistent with recent findings by Hu et al. (2025), who noted that Chinese students' attitudes toward AI are closely associated with direct classroom exposure and teacher guidance. According to SCT, supervisors' ethical standards may shape students' ethical awareness and behavior through modeling, guidance, and repeated academic interaction (Nejati & Shafaei, 2018). While an academic ethical climate remains important, it may be too abstract to be directly reflected in students' self-reported ethical LLM use, and it may require time to be internalized into students' values and beliefs (Kelman, 1958; Cheng et al., 2021). Therefore, universities should not rely only on broad institutional declarations. They also need supervisor-level guidance, concrete examples, and practical support to help students understand and apply responsible AI-use standards.

Besides, ethical supervision requires that students, teachers, administrators, and marginalized community members co-design AI guidelines. This ensures that the governance framework reflects the diverse pedagogical needs and cultural contexts of the entire school ecosystem. By treating AI integration as a social challenge rather than a purely technological one, institutions can establish ethical supervision mechanisms that are transparent, structurally inclusive, and adaptable to the evolving landscape of higher education.

The Ambiguity of Internal Cognition: Moral Consciousness and Self-Efficacy

While moral consciousness and ethical chatbot self-efficacy were positively associated with ethical use, the strength of these relationships was weaker than anticipated by traditional ethical decision-making literature (Pan & Sparks, 2012; Wang et al., 2013).

This weakened link can be explained by the current low social consensus surrounding LLM usage ethics (Barnett & Valentine, 2004; Zhang et al., 2025). The rapid evolution of these tools raises complex questions about accountability, data privacy, and transparency, leaving students without a clear moral compass. Consequently, as Zhang et al. (2025) observed, ambiguous social norms often lead to the “secret use” of LLMs rather than principled ethical engagement. The novelty of LLMs means students lack established ethical frameworks, having missed formative moral education regarding these specific technologies.

Similarly, ethical chatbot self-efficacy—reflecting students’ confidence in controlling their behavior—may be undermined by environmental anxiety. As self-efficacy perceptions influence emotional arousal (Bandura, 1997), students may feel their ethical actions are disadvantaged if they observe peers misusing LLMs to gain an unfair edge without facing consequences. This peer-induced stress, compounded by unclear institutional guidelines, reduces their confidence in maintaining ethical behavior (Farrell et al., 2023; Malesky et al., 2022).

When ethical guidelines for AI use are communicated transparently and implemented consistently, AI governance shifts from a predominantly top-down enforcement approach toward a shared framework of institutional norms and collective responsibility. Such clarity enables teachers to design assessments that appropriately integrate AI technologies, helps students understand the ethical boundaries governing LLM use, and equips guidance counselors to recognize and address the potential effects of AI-related stress on student well-being. By fostering a transparent and supportive educational environment, these governance practices are expected to strengthen students’ moral consciousness and self-efficacy, thereby reinforcing their positive associations with ethical LLM use.

The Paradox of Policy Awareness and Prolonged Use

A notable finding is that awareness of chatbot usage policies was not significantly associated with ethical use, whereas longer experience with chatbot use was negatively associated with ethical behavior. Unlike the findings of Mvondo et al. (2024), this result suggests that policy awareness alone may not be sufficient to guide students’ ethical LLM use. One possible explanation is that students may perceive current AI detection and policy enforcement mechanisms as limited or inconsistent (Androćec, 2024; Hong, 2026). Although length of chatbot use was included only as a post-hoc control variable, its negative association with ethical use may reflect students’ increasing familiarity with both the efficiency benefits of LLMs and the limitations of institutional monitoring. From an RCT perspective, students with longer experience may become more aware of how LLMs can reduce workload, while also recognizing possible ways to avoid detection or accountability (Androćec, 2024; Sadasivan et al., 2023). However, this interpretation should remain cautious because the study design does not establish causality. Prolonged exposure to LLMs may also be associated with greater normalization of AI-assisted shortcuts, such as copying AI-generated outputs without sufficient understanding or using fabricated references without verification (Hong & Litwin, 2025; Yan et al., 2025). From an SCT perspective, peer practices may further shape students’ perceptions of what counts as acceptable AI use, especially when inappropriate LLM use appears common or goes unchallenged (Bandura, 1986; Hong et al., 2023).

Nevertheless, given the cross-sectional nature of the data, this negative association should not be interpreted as evidence that longer chatbot use causes less ethical behavior. Several alternative explanations are possible. For instance, novice users may use LLMs in fewer academic situations, whereas experienced users may apply them across a wider range of tasks, increasing both opportunities and temptations for inappropriate use. Additionally, students with longer exposure to LLMs may have greater awareness of what counts as unethical use, while novice users may overestimate their own ethical compliance due to limited AI literacy (An et al., 2023). Therefore, while the RCT perspective offers a plausible explanation, the relationship should be understood as associational and exploratory, requiring further longitudinal investigation.

External Deterrence and the SCT-RCT Interplay

The study confirms that both AI writing detection probability and sanction severity are positively related to ethical use, consistent with deterrence theory (Wenzel, 2004). Under RCT, this reflects a rational calculation: students are more likely to behave ethically when they perceive a high likelihood of being caught and severely punished. However, these connections remain somewhat weak, likely because, as stated above, students are skeptical about the actual effectiveness of current AI detection tools, which are notoriously unreliable and prone to bias against non-native speakers (Hong, 2026). Furthermore, an over-reliance on sanctions risks creating a counter-productive “policing culture” that damages relational trust between students and faculty (Song, 2024).

Interestingly, the parallel effects of internal moral consciousness and external detection possibilities suggest a balanced interplay between SCT and RCT. Students operate within a complex decision-making framework guided by both moral reasoning and rational risk assessment. These findings highlight the importance of adopting a holistic institutional strategy that integrates ethical supervision (to build moral consciousness via SCT) with credible and transparent deterrence mechanisms (to influence rational calculations via RCT) to promote the responsible use of LLMs.

Practical Implications

Suggestions for Higher Education Institutions

Effectively integrating LLMs into academic practices rather than implementing outright bans (Hong, 2023) remains the most important subject for higher education institutes. To balance innovation with integrity, institutions should:

Particularly with the rapid rollout of highly capable domestic models like Deepseek and Wenxin Yiyan (Neha & Bhati, 2025; Tian et al., 2024), it is imperative that institutions develop comprehensive regulatory frameworks with clear boundaries. A prime example is Prague University’s 2024 AI framework, which explicitly delineates acceptable uses—such as autodidactic purposes and stylistic control—from unacceptable ones, thereby eliminating ambiguity and providing a firm foundation for ethical awareness (Dobrovská et al., 2024). Building upon these clear regulations, universities must also establish balanced oversight. By adopting frameworks similar to the European Union’s Artificial Intelligence Act (European Commission, 2021), institutions should implement minimum standards for transparency and attribution through a multi-layered governance approach, which helps avoid the pitfalls of a disproportionately punitive “policing culture” (Hong, 2026). Furthermore, the success of such oversight relies heavily on the presence of authoritative ethical role models. Given that authority figures hold significant influence, especially in East Asian cultures (Ghosh & Chatterjee, 2024), university leaders and faculty must actively demonstrate ethical AI use to establish a unified standard and strengthen the social value of academic integrity. Ultimately, to ensure these structural and cultural shifts take root, institutions must systematically strengthen ethical awareness and education. Although many Chinese universities currently offer online AI courses, launching dedicated AI ethics training within the core curriculum remains an urgent necessity to help students internalize responsible practices and combat the increasingly blurred boundaries of originality (Hutson, 2024; Song, 2024).

Suggestions for Students

From the student’s perspective, balancing the productivity benefits of LLMs with academic integrity requires proactive self-regulation:

To effectively navigate the integration of LLMs in academic settings, students must actively cultivate moral self-efficacy by building confidence in their ability to use these tools responsibly and transparently. Instructors can facilitate this process by creating opportunities for students to openly disclose their AI usage and collaboratively identify the technology’s limitations. Furthermore, forming peer “supervision-sharing” groups can help establish positive social norms, thereby counteracting the anxiety associated with unfair peer advantages. Alongside developing this ethical confidence, students must also consciously engage with institutional regulations and standardize their citations. By utilizing citation plugins to accurately annotate AI assistance, they can avoid the pitfalls of a “quick-fix” mentality that often leads to cognitive atrophy (Hong, 2026). Ultimately, these individual practices are essential for fostering a positive usage environment. By publicly committing to academic integrity and forming interest groups focused on responsible AI use, students can effectively leverage the “bandwagon effect” from SCT to influence collective norms in a positive and sustainable direction.

Limitations and Future Directions

This study has several limitations that present opportunities for future research. First, the sample primarily consisted of undergraduate students under 25 years old, heavily concentrated in the social sciences and engineering. Future studies should adopt stratified sampling to capture discipline-specific differences—especially given that engineering students exhibit distinct, high-frequency usage patterns (Sun et al., 2025)—and include mature or post-graduate students, whose closer supervisory relationships may yield different structural patterns. Second, because the study relies on cross-sectional, self-reported data, we cannot establish causal relationships between the variables. All findings represent structural associations, and self-reported measures may carry inherent social desirability biases. Longitudinal studies are needed to track how prolonged exposure relates to changes in ethical behavior and self-regulation over time, and to rule out unmeasured institutional factors. Third, the study exclusively examines student perspectives, omitting faculty and administrators whose insights would provide a more holistic understanding of AI governance. Finally, although this study did not aim to validate the scales, initial psychometric analyses revealed that Cronbach's alphas were occasionally higher than composite reliability (CR) estimates, suggesting possible item redundancy. While rigorous sensitivity and split-half analyses confirmed overall construct reliability without compromising the model, future research should replicate this adapted scale to further validate its psychometric properties. Lastly, our reliance on a single-item measure for Ethical Use, although referencing existing literature and to minimize respondent defensiveness, future studies could adopt multidimensional scales to measure ethical AI behavior.

Conclusions

Nevertheless, given the cross-sectional nature of the data, this negative association should not be interpreted as evidence that longer chatbot use causes less ethical behavior. Several alternative explanations are possible. For instance, novice users may use LLMs in fewer academic situations, whereas experienced users may apply them across a wider range of tasks, increasing both opportunities and temptations for inappropriate use. Additionally, students with longer exposure to LLMs may have greater awareness of what counts as unethical use, while novice users may overestimate their own ethical compliance due to limited AI literacy (An et al., 2023). Therefore, while the RCT perspective offers a plausible explanation, the relationship should be understood as associational and exploratory, requiring further longitudinal investigation.

In practice, the findings suggest that universities should treat ethical AI use as both a pedagogical and governance issue, rather than only a technical concern (Floridi et al., 2018). Passive publication of AI-use policies may be insufficient to guide students' behavior. Instead, institutions should combine clear policies with supervisor-level guidance, authentic assessment design, and fair accountability mechanisms. Students need concrete examples, transparent expectations, and opportunities to learn how to use LLMs responsibly in academic tasks. By addressing both ethical guidance and external accountability, universities can better support responsible LLM use while maintaining academic integrity, student trust, and inclusive learning environments.

Author Contributions

CZ is the first author, taking charge in conceptualization, writing the original draft, reviewing and editing. WCHH is the second author, in charge of methodology, software, formal analysis and editing. NC is the third author, reviewing. YFZ is the fourth author as well as co-corresponding author, doing supervision. XSX is the corresponding author, project administration and funding acquisition. All authors read and approved the final manuscript. All authors agree to be accountable for all aspects of the work in ensuring that questions related to the accuracy or integrity of any part of the work are appropriately investigated and resolved.

Ethics Approval and Consent to Participate

The study received ethical approval from the Institutional Review Board (IRB) of Wenzhou University. IRB LOG # (Assigned by IRB) #20260226. Anticipated starting and completion dates: 02/28/2026 to 03/27/2026.

Funding

This paper is supported by the Macao Polytechnic Research Fund (Grant#: RP/FLT-06/2025).

Availability of Data and Materials

Data can be accessed from: <https://figshare.com/s/9b471bd003623fe08b82>. Doi: 10.6084/m9.figshare.29502140.

Conflict of Interest

The author(s) declared no potential conflicts of interest with respect to the research, authorship, and/or publication of this article.

Acknowledgments

Not applicable.

Appendix

Exploratory Factor Analysis

A principal component analysis was conducted on 33 items with varimax rotation (Table 8). The Kaiser–Meyer–Olkin measure verified the sampling adequacy, at .96. Bartlett’s test of sphericity, $\chi^2(528) = 16,893.47$, $p < .001$, showed that correlations between items were sufficiently large for Principal Component Analysis.

Using Varimax rotation and Kaiser Normalization. Results converged in 7 iterations, with all items above .5 in factor loading, confirming high construct validity in each component. Exactly as Mvondo et al. (2024) reported, EFA extracted 7 components explaining 79.15% of the total variance, with the main component accounting for 51.7% variance. The average variance extracted was at or above the .5 recommended threshold, except for Awareness of chatbot usage policies, at .49. Composite reliability all exceeded .7 (See Table 2). VIF was below 5. Taken together, construct validity and reliability were confirmed, while multicollinearity was not a problem.

Effects of Other Factors

A series of multiple regression analyses were conducted to determine if demographic and other factors—Age, Gender, Academic status, Discipline and Length of chatbot usage—had direct or mediating effects on Ethical use. Results confirmed several significant predictors of EU, Ethical supervision, $\beta = .21$, $p = .001$, 95% CI [.10, .42]. Moral consciousness, $\beta = .19$, $p = .002$, 95% CI [.09, .37]. Ethical chatbot self-efficacy, $\beta = .15$, $p = .009$, 95% CI [.04, .26], AI writing detection probability $\beta = .11$, $p = .006$, 95% CI [.08, .48] and Sanction severity, $\beta = .09$, $p = .037$, 95% CI [.01, .20]. Interestingly, Length of chatbot use (Length) had a negative effect on EU, $\beta = -.32$, $p < .001$, 95% CI [−.19, −.12].

For the moderators, we found Length had a positive effect on AI usage, $\beta = .06$, $p = .036$, 95% CI [.00, .06]. Thus, it was included in the Structural Equation Model.

Table 8. Construct validity and reliability results.

Construct items	Loading	AVE	CR
Ethical climate		.61	.88
EC1: My university has formal guidelines on the ethical use of LLM chatbots.	.78		
EC2: My university strictly enforces academic integrity.	.75		
EC3: My university has policies regarding the ethical use of LLM chatbots.	.84		
EC4: My university strictly enforces policies regarding the ethical use of LLM chatbots.	.81		
EC5: Professors in my university have let it be known in no uncertain terms that unethical use of LLM chatbots will not be tolerated.	.73		
Ethical supervision		.50	.90
ES1: My supervisor listens to what students have to say.	.62		
ES2: My supervisor disciplines students who violate ethical standards.	.51		
ES3: My supervisor conducts his/her personal life ethically.	.67		
ES4: My supervisor has the best interests of students in mind.	.77		
ES5: My supervisor makes fair and balanced decisions.	.80		
ES6: My supervisor can be trusted.	.83		
ES7: My supervisor discusses academic ethics or values with students.	.63		
ES8: My supervisor sets an example of how to do things correctly in terms of ethics.	.78		
ES9: My supervisor defines success not just by results but also by the way they are obtained.	.68		
ES10: When making decisions, my supervisor asks, "What is the right thing to do?"	.66		
Moral consciousness		.50	.86
In using LLM chatbots,			
MC1: I often have to choose between doing what is right and doing something wrong.	.67		
MC2: I regularly face decisions that have significant ethical implications.	.73		
MC3: Many of the decisions that I make have ethical dimensions to them.	.76		
MC4: I regularly think about the ethical implications of my decisions.	.77		
MC5: I think about the morality of my actions almost every day.	.69		
MC6: I like to think about ethics.	.61		
Ethical chatbot self-efficacy		.66	.91
ECSE1: When you need to write an essay or a research article but feel that you cannot do it yourself, how confident are you to refuse to use a chatbot unethically?	.76		
ECSE2: When you need to write an essay or a research article and the deadline is fast approaching, how confident are you to refuse to use a chatbot unethically?	.80		
ECSE3: When you need to write an essay or a research paper, and you can use a chatbot, how confident are you not to take advantage of it?	.85		
ECSE4: When you need to write an essay or a research article and have seen other students using a chatbot for this purpose, how confident are you not to take advantage of it?	.87		
ECSE5: When you need to use a chatbot for your essay or research paper, how confident are you not to take advantage of it?	.80		
Awareness of chatbot usage policies		.49	.74
ACUP1: I am aware of the usage policies prohibiting using LLM chatbots to engage in or promote academic dishonesty.	.67		
ACUP2: I understand my responsibilities as outlined in the AI chatbots' usage policies.	.75		
ACUP3: I comprehend the policies set forth by AI chatbot companies for the ethical use of chatbots.	.67		
AI writing detection probability		.77	.87
AIWDP1: If I used LLM chatbots unethically, the probability that it would be detected is (Very Low/Very High).	.89		
AIWDP2: If I used LLM chatbots unethically, I would probably be caught.	.87		
Sanction severity		.70	.82
SS1: I think the punishment would be (Very Low/Very High).	.84		
SS2: I would be severely punished by my university.	.83		

References

- An, Q., Hong, W. C. H., Xu, X. S., Zhang, Y., & Kolletar-Zhu, K. (2023). How education level influences internet security knowledge, behaviour, and attitude: A comparison among undergraduates, postgraduates and working graduates. *International Journal of Information Security*, 22(2), 305–317. <https://doi.org/10.1007/s10207-022-00637-z>
- Androćec, D. (2024). Using large language models for students' essays plagiarism detection. In *2024 5th International Conference on Communications, Information, Electronic and Energy Systems (CIEES)* (pp. 1–5). IEEE. <https://doi.org/10.1109/CIEES62939.2024.10811263>
- Bandura, A. (1986). *Social foundations of thought and action: A social cognitive theory*. Prentice Hall.
- Bandura, A. (1997). *Self-efficacy: The exercise of control*. W. H. Freeman.
- Barnett, T., & Valentine, S. (2004). Issue contingencies and marketers' recognition of ethical issues, ethical judgments and behavioral intentions. *Journal of Business Research*, 57(4), 338–346. [https://doi.org/10.1016/S0148-2963\(02\)00365-X](https://doi.org/10.1016/S0148-2963(02)00365-X)
- Barrot, J. S. (2023). Using ChatGPT for second language writing: Pitfalls and potentials. *Assessing Writing*, 57, 100745. <https://doi.org/10.1016/j.asw.2023.100745>
- Bass, D. (2022, December 8). *OpenAI chatbot so good it can fool humans, even when it's wrong*. Bloomberg. <https://www.bloomberg.com/news/articles/2022-12-07/openai-chatbot-so-good-it-can-fool-humans-even-when-it-s-wrong>
- Becker, G. S. (1968). Crime and punishment: An economic approach. *Journal of Political Economy*, 76(2), 169–217. <https://www.jstor.org/stable/1830482>
- Cheng, Y. C., Hung, F. C., & Hsu, H. M. (2021). The relationship between academic dishonesty, ethical attitude and ethical climate: The evidence from Taiwan. *Sustainability*, 13(21), 11615. <https://doi.org/10.3390/su132111615>
- Chinese University of Hong Kong. (2023). *Use of artificial intelligence tools in teaching, learning, and assessment: A guide for students*. https://www.aqs.cuhk.edu.hk/documents/A-Guide-for-students_use-of-AI-tools.pdf
- Cohen, J. (1988). *Statistical power analysis for the behavioral sciences (2nd ed.)*. Routledge. <https://doi.org/10.4324/9780203771587>
- Dergaa, I., Chamari, K., Zmijewski, P., & Ben Saad, H. (2023). From human writing to artificial intelligence-generated text: Examining the prospects and potential threats of ChatGPT in academic writing. *Biology of Sport*, 40(2), 615–622. <https://doi.org/10.5114/biolsport.2023.125623>
- Dobrovská, D., Vaněček, D., & Yorulmaz, Y. I. (2024). Ethical rules in using AI for educational purposes: A case study of CTU in Prague. In *International Conference on Interactive Collaborative Learning* (pp. 368–376). Springer. https://doi.org/10.1007/978-3-031-85652-5_37
- Dowling, M., & Lucey, B. (2023). ChatGPT for (finance) research: The Bananarama conjecture. *Finance Research Letters*, 53, 103662. <https://doi.org/10.1016/j.frl.2023.103662>
- Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., Baabdullah, A. M., Koohang, A., Raghavan, V., Ahuja, M., Albanna, H., Albashrawi, M. A., Al-Busaidi, A. S., Balakrishnan, J., Barlette, Y., Basu, S., Bose, I., Brooks, L., Buhalis, D., ... Wright, R. (2023). “So what if ChatGPT wrote it?” Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *International Journal of Information Management*, 71, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- European Commission. (2021). *Proposal for a regulation laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) and amending certain Union legislative acts*. EUR-Lex. <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=celex%3A52021PC0206>
- Farrell, W. C., Bogodistov, Y., & Mössenlechner, C. (2023). Is academic integrity at risk? Perceived ethics and technology acceptance of ChatGPT. *AMCIS 2023 Proceedings*. Association for Information Systems. https://aisel.aisnet.org/amcis2023/sig_ed/sig_ed/50/
- Farrokhnia, M., Banihashem, S. K., Noroozi, O., & Wals, A. (2024). A SWOT analysis of ChatGPT: Implications for educational practice and research. *Innovations in Education and Teaching International*, 61(3), 460–474. <https://doi.org/10.1080/14703297.2023.2195846>
- Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P., & Vayena, E. (2018). AI4People—An ethical framework for a good AI society: Opportunities, risks, principles, and recommendations. *Minds and Machines*, 28, 689–707. <https://doi.org/10.1007/s11023-018-9482-5>
- Gallent-Torres, C., Zapata-González, A., & Ortego-Hernando, J. L. (2023). The impact of generative artificial intelligence in higher education: A focus on ethics and academic integrity. *RELIEVE*, 29(2), 1–19. <https://doi.org/10.30827/relieve.v29i2.29134>
- Ghosh, S., & Chatterjee, S. (2024). Machine translation, large language models, and generative AI in the university classroom: Toward a pedagogy of care. In *The social impact of automating translation* (pp. 163–175). Routledge. <https://doi.org/10.4324/9781003465522-9>
- Gorelick, E., & McDonald, A. (2023, February 12). *University leaders issue AI guidance in response to growing popularity of ChatGPT*. Yale Daily News. <https://yaledailynews.com/blog/2023/02/12/university-leaders-issue-ai-guidance-in-response-to-growing-popularity-of-chatgpt/>
- Hair, J. F., Hult, G. T. M., Ringle, C. M., & Sarstedt, M. (2022). *A primer on partial least squares structural equation modeling (PLS-SEM)* (3rd ed.). Sage.
- Hair, J. F., Risher, J. J., Sarstedt, M., & Ringle, C. M. (2019). When to use and how to report the results of PLS-SEM. *European Business Review*, 31(1), 2–24. <https://doi.org/10.1108/EBR-11-2018-0203>
- Henseler, J., Hubona, G., & Ray, P. A. (2016). Using PLS path modeling in new technology research: Updated guidelines. *Industrial Management & Data Systems*, 116(1), 2–20. <https://doi.org/10.1108/IMDS-09-2015-0382>
- Hofstede, G. (2011). Dimensionalizing cultures: The Hofstede model in context. *Online Readings in Psychology and Culture*, 2(1), 8. <https://doi.org/10.9707/2307-0919.1014>
- Hong, W. C. H. (2023). The impact of ChatGPT on foreign language teaching and learning: Opportunities in education and research. *Journal of Educational Technology and Innovation*, 5(1), 37–45. <https://doi.org/10.61414/jeti.v5i1.103>

- Hong, W. C. H. (2026). Reflecting on three years of generative AI in higher education: Successes, struggles, and solutions. *Artificial Intelligence Advances in Education*, 1(1), 1–11. <https://doi.org/10.66818/aiaie.v1i1.916>
- Hong, W. C. H., & Litwin, A. (2025). From human guidance to digital assistance: EFL thesis-writing students' experiences of support before and after the AI revolution. *Journal of Educational Technology and Innovation*, 7(3), 21–34. <https://doi.org/10.61414/a3wy7546>
- Hong, W. C. H., Chi, C., Liu, J., Zhang, Y., Lei, V. N. L., & Xu, X. (2023). The influence of social education level on cybersecurity awareness and behaviour: A comparative study of university students and working graduates. *Education and Information Technologies*, 28(1), 439–470. <https://doi.org/10.1007/s10639-022-11121-5>
- Hu, X., Gong, W., & Cortazzi, M. (2025). Researching China's pre-service teachers' perceptions of AI education for K–12: An elicited metaphor analysis. *European Journal of Education*, 60(2), e70079. <https://doi.org/10.1111/ejed.70079>
- Hutson, J. (2024). Rethinking plagiarism in the era of generative AI. *Journal of Intelligent Communication*, 3(2), 20–31. <https://doi.org/10.54963/jic.v3i2.220>
- Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., Gasser, U., Groh, G., Günemann, S., Hüllermeier, E., Krusche, S., Kutyniok, G., Michaeli, T., Nerdel, C., Pfeffer, J., Poquet, O., Sailer, M., Schmidt, A., Seidel, T., ... Kasneci, G. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, 102274. <https://doi.org/10.1016/j.lindif.2023.102274>
- Kelman, H. C. (1958). Compliance, identification, and internalization: Three processes of attitude change. *Journal of Conflict Resolution*, 2(1), 51–60. <https://doi.org/10.1177/002200275800200106>
- Kong, S. C., & Zhu, J. (2025). Developing and validating an artificial intelligence ethical awareness scale for secondary and university students: Cultivating ethical awareness through problem-solving with artificial intelligence tools. *Computers and Education: Artificial Intelligence*, 9, 100447. <https://doi.org/10.1016/j.caeai.2025.100447>
- Kung, T. H., Cheatham, M., Medenilla, A., Sillos, A., De Leon, L., Elepaño, C., Madriaga, M., Aggabao, R., Diaz-Candido, G., Maningo, J., & Tseng, V. (2023). Performance of ChatGPT on USMLE: Potential for AI-assisted medical education using large language models. *PLOS Digital Health*, 2(2), e0000198. <https://doi.org/10.1371/journal.pdig.0000198>
- Li, M., Chen, Y.-T., Huang, C.-Q., Hwang, G.-J., & Cukurova, M. (2023). From motivational experience to creative writing: A motivational AR-based learning approach to promoting Chinese writing performance and positive writing behaviours. *Computers & Education*, 202, 1–23. <https://doi.org/10.1016/j.compedu.2023.104844>
- Li, S., Huang, J., Hussain, S., & Dong, Y. (2025). How does supervisor support impact Chinese graduate students' research creativity through research self-efficacy and intrinsic motivation? A multi-group analysis. *Thinking Skills and Creativity*, 55, 101700. <https://doi.org/10.1016/j.tsc.2024.101700>
- Lv, Z. (2023). Generative artificial intelligence in the metaverse era. *Cognitive Robotics*, 3, 208–217. <https://doi.org/10.1016/j.cogr.2023.06.001>
- Malesky, A., Grist, C., Poovey, K., & Dennis, N. (2022). The effects of peer influence, honor codes, and personality traits on cheating behavior in a university setting. *Ethics & Behavior*, 32(1), 12–21. <https://doi.org/10.1080/10508422.2020.1869006>
- Martin, K. D., & Cullen, J. B. (2006). Continuities and extensions of ethical climate theory: A meta-analytic review. *Journal of Business Ethics*, 69(2), 175–194. <https://doi.org/10.1007/s10551-006-9084-7>
- Microsoft Education. (2025). *2025 AI in education: A Microsoft special report*. Microsoft. <https://cdn-dynmedia-1.microsoft.com/is/content/microsoftcorp/microsoft/bade/documents/products-and-services/en-us/education/2025-Microsoft-AI-in-Education-Report.pdf>
- Mvondo, G. F. N., Niu, B., & Eivazinezhad, S. (2024). Exploring the ethical use of LLM chatbots in higher education. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.4548263>
- Neha, F., & Bhati, D. (2025). *A survey of DeepSeek models*. TechRxiv. <https://doi.org/10.36227/techrxiv.173896582.25938392.v1>
- Nejati, M., & Shafaei, A. (2018). Leading by example: The influence of ethical supervision on students' prosocial behavior. *Higher Education*, 75(1), 75–89. <https://doi.org/10.1007/s10734-017-0130-4>
- Pan, Y., & Sparks, J. R. (2012). Predictors, consequence, and measurement of ethical judgments: Review and meta-analysis. *Journal of Business Research*, 65(1), 84–91. <https://doi.org/10.1016/j.jbusres.2011.02.002>
- Rout, K., Sahoo, P. R., Behera, P., & Bhuyan, A. (2024). Cognitive, affective, and demographic factors and their impact on human behavior to adopt LLM-based generative artificial intelligence. In *Future of management: Embracing sustainability, diversity, and inclusivity* (pp. 71–82). Routledge.
- Sadasivan, V. S., Kumar, A., Balasubramanian, S., Wang, W., & Feizi, S. (2023). *Can AI-generated text be reliably detected?* arXiv. <https://doi.org/10.48550/arXiv.2303.11156>
- Schei, O. M., Møgelvang, A., & Ludvigsen, K. (2024). Perceptions and use of AI chatbots among students in higher education: A scoping review of empirical studies. *Education Sciences*, 14(8), 922. <https://doi.org/10.3390/educsci14080922>
- Seidman, A. (2019). *Minority student retention: The best of the "journal of college student retention: research, theory & practice"*. Routledge.
- Shahzad, M. F., Xu, S., Lim, W. M., Yang, X., & Khan, Q. R. (2024). Artificial intelligence and social media on academic performance and mental well-being: Student perceptions of positive impact in the age of smart learning. *Heliyon*, 10(8), e29523. <https://doi.org/10.1016/j.heliyon.2024.e29523>
- Song, N. (2024). Higher education crisis: Academic misconduct with generative AI. *Journal of Contingencies and Crisis Management*, 32(1), e12532. <https://doi.org/10.1111/1468-5973.12532>
- Sun, Y., Yang, H., Yu, H. K., & Suen, R. (2025). Boon or bane? Evaluating AI-driven learning assistance in higher education professional coursework.

- Education and Information Technologies*, 30, 22011–22044. <https://doi.org/10.1007/s10639-025-13642-1>
- Susnjak, T. (2022). *ChatGPT: The end of online exam integrity?* arXiv. <https://doi.org/10.48550/arXiv.2212.09292>
- Tian, W., Ge, J., Zhao, Y., & Zheng, X. (2024). AI chatbots in Chinese higher education: Adoption, perception, and influence among graduate students—An integrated analysis utilizing UTAUT and ECM models. *Frontiers in Psychology*, 15, 1268549. <https://doi.org/10.3389/fpsyg.2024.1268549>
- University of California, Los Angeles. (2023). *Guidance for the use of generative AI*. https://teaching.ucla.edu/resources/ai_guidance/#toggle-id-7
- Wang, Y.-S., Yeh, C.-H., & Liao, Y.-W. (2013). What drives purchase intention in the context of online content services? The moderating role of ethical self-efficacy for online piracy. *International Journal of Information Management*, 33(1), 199–208. <https://doi.org/10.1016/j.ijinfomgt.2012.09.004>
- Wenzel, M. (2004). The social side of sanctions: Personal and social norms as moderators of deterrence. *Law and Human Behavior*, 28(5), 547–567. <https://doi.org/10.1023/B:LAHU.0000046433.57588.71>
- Yan, X., Zhu, K., Yang, G., Liu, T., Chen, F., & Lu, Y. (2025). The influence and analysis of large language model on college students' homework and other course tasks. *Proceedings of the 2025 International Conference on Education Research and Training Technologies (ERTT 2025)* (pp. 221–229). Atlantis Press. https://doi.org/10.2991/978-2-38476-479-2_26
- Zhang, Z., Shen, C., Yao, B., Wang, D., & Li, T. (2025). Secret use of large language model (LLM). *Proceedings of the ACM on Human-Computer Interaction*, 9(2), 1–26. <https://doi.org/10.1145/3711061>