

International Journal of Learning, Teaching and Educational Research
 Vol. 25, No. 8, pp. 641-660, August 2026
<https://doi.org/10.26803/ijlter.25.8.25>
 Received May 10, 2026; Revised Jul 8, 2026; Accepted Jul 10, 2026

From Innovation to Imitation: Associations Between Exposure to AI Prompt-Sharing Content on Social Media, Moral Disengagement, and Intention to Engage in Academic Misconduct

Yuansheng Liu* 

Universiti Putra Malaysia
 Seri Kembangan, Selangor, Malaysia

Xinjue Jiang 

Universiti Malaya
 Kuala Lumpur, Malaysia

Abstract. The use of AI today has raised significant ethical concerns regarding its potential misuse in academic contexts. This paper aims to examine relationship between university students' exposure to AI prompt-sharing content on social media, moral disengagement, perceived academic norms, and intentions to engage in AI-assisted academic misconduct. Specifically, the study investigated whether exposure to AI prompt-sharing content was associated with students' intentions to engage in AI-assisted academic misconduct, whether moral disengagement mediated this relationship, and whether perceived academic norms moderated the association between exposure and moral disengagement. The study conducted SPSS-based descriptive statistics, regression, mediation, and moderation analyses of data collected through a structured questionnaire administered to 600 Chinese university students to verify three hypotheses. The findings indicate that greater exposure to AI prompt-sharing content was significantly associated with stronger intentions to engage in AI-assisted academic misconduct, with moral disengagement acting as a significant mediator, and perceived academic norms strengthening the association between exposure and moral disengagement. These observations indicate that AI prompt-sharing content on social media can be associated with students' ethical decision-making regarding AI use in academic settings and highlight the importance of promoting responsible AI use and academic integrity in higher education.

Keywords: AI prompt-sharing; academic misconduct; moral disengagement; social media

Citation :

Liu, Y. S., & Jiang, X. J.
 (2026) From Innovation to Imitation: Associations Between Exposure to AI Prompt-Sharing Content on Social Media, Moral Disengagement, and Intention to Engage in Academic Misconduct. *International Journal of Learning, Teaching and Educational Research*, 25(8), 641-660.
<https://doi.org/10.26803/ijlter.25.8.25>

*Corresponding author: Yuansheng Liu; 218480@student.upm.edu.my

1. Introduction

Social media platforms make possible a significant degree of communication on various subjects by users. Social media is a new mode of communication that influences the way its users, all over the world, interact with the platforms to obtain information; in the process their perceptions are shaped by their interaction with individuals and organizations (Alalwan et al., 2017; Appel et al., 2020; Lee, 2022). Kotler et al. (2022) state that social media has become a way for users to share information, in the form of images, text, videos, and audio, with others. Such communication influences users' approaches to subject matter, and they are guided by their search patterns and by algorithmic curation by the platform. An algorithmic curation-based search promotes participation and influences the internalization of norms and standards by users.

This algorithmic curation process is described by Narayanan (2023), who reports that social media platforms predict the patterns of searching and content engagement of users; the platform promotes content that generates the most user responses in the form of likes, shares and comments. These trends enable the platforms to curate content of a similar pattern through a repeated process that increases the visibility of a certain category of content and the exposure of users to such categories. Platforms rank posts in order of decreasing prediction of engagement. For example, Facebook suggests posts that users are most likely to like, react to, share, or comment on; YouTube suggests videos that are most likely to gain clicks, and those that are likely to have longer watch time (Narayanan, 2023). This algorithmic recommendation of content based on predicted user engagement applies to content of any category, including content that promotes academic strategies guiding students on various approaches to performance optimization and increasing academic productivity.

Students' academic performance is defined and shaped by universities and their administrative policies, which inform students about acceptable behaviors, guidelines and what constitutes misconduct. For university students, academic integrity includes avoiding plagiarism in their academic tasks, not cheating, and finishing their tasks independently (Wong et al., 2016). Amigud and Pell (2022) report that staff of academic institutions use improvised strategies such as doing selective screening, issuing informal warnings, and reframing messages as measures to curb integrity violations, while balancing institutional expectations with personal values.

They note that standardization of policies and fair practices is hampered by context-specific decision-making that shows variation, and inconsistency in integrity enforcement that is guided by personal experience. Furthermore, academic misconduct is related to institutional influence. Ahmad et al. (2025) reports that students indulge in such practices because of an absence of instructional guidance on plagiarism, the individual desire to achieve good grades, and students' fear of failure. University administrations exert considerable pressure on students to perform well despite not providing significant measures to curb students' engagement in misconduct related to integrity.

Popescu et al. (2023) suggest that AI can be useful and reduce costs when it is integrated within an administrative framework to perform repetitive work that does not require human intellectual and manual contributions. While such functions ensure the positive integration of AI in education practices, it does not suggest ways to control academic misconduct by students. Furthermore, the literature does not discuss the influence of external social pressures on the motivation of the students to indulge in academic misconduct.

1.1 Research Aim, Questions, and Hypotheses

This study investigated the way social media communication shapes AI use norms in academic contexts of Chinese students. The core purpose of this study was to identify how social media communication becomes a mechanism for diffusing the established norms of AI use in academia. Therefore, the purpose of this research was to investigate the relationship between exposure to content on AI prompt sharing and student motivation for academic misconduct through unethical AI use in academia in constructing a collective perception among the students of peer norms and whether exposure to such AI-content sharing triggers moral-cognitive processes such as moral disengagement, which cause students to rationalize unethical use of AI in academic practices.

Consequently, by investigating the patterns through which social media normalizes AI-assisted shortcuts in academia, the paper investigates the socio-psychological pathways that promote the cultural adoption and legitimization of AI-based academic misconduct. In this research study, AI use refers to the overall use of generative AI tools; AI prompt sharing refers to the process of sharing prompts through social media; and academic misconduct through the use of AI refers to the unethical use of AI for unfair academic gain. This research addressed two research questions:

1. Do AI prompt-sharing videos influence perceived academic norms?
2. Does moral disengagement mediate the link between exposure and misconduct intention?

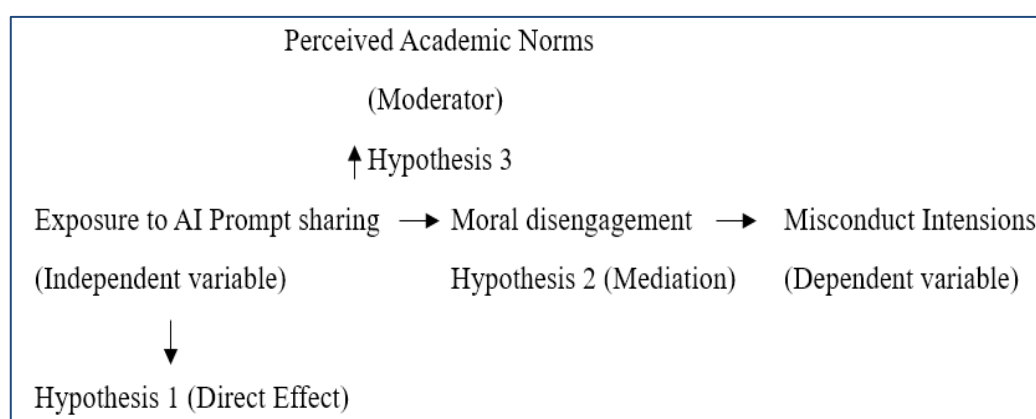


Figure 1: Diagrammatic representation of the correlation between independent variable, dependent variable, and hypotheses

To address these research questions, the study established three hypotheses that operated through their correlation with the independent and dependent variables of the study, and which were intended to be tested and verified in this study:

H₁: Exposure to AI prompt-sharing content via social media is positively associated with students' intentions to engage in AI-assisted academic misconduct.

H₂: Moral disengagement mediates the relationship between exposure to AI prompt-sharing content and intentions of academic misconduct.

H₃: Perceived academic norms moderate the relationship between exposure to AI prompt-sharing content and moral disengagement, such that the relationship is stronger when norms are more permissive.

1.2 Significance of Study

The significance of this research is that it offers an empirical evaluation of how social media posts on AI prompt sharing influence the academic use of AI through exposing students to such posts. The paper reveals how exposure to AI prompt-sharing posts on social media normalizes academic misconduct for students by developing communication pathways that communally diffuse unethical practices in academia. The paper portrays the psychological mechanism of moral disengagement as the primary linkage between online behavior and students' intention to engage in academic misconduct.

The paper identifies academic norms as a social influence tool that contributes to the restraint or amplification of unethical use of AI by students. This research examines the gap identified in the literature on the correlation between exposure to AI prompt-sharing content on social platforms, students' perceptions of academic norms, their moral disengagement, and their intention to engage in behaviors involving academic dishonesty. This study is significant in the current time because of the increased use of AI in everyday domains, which leads to questions about the distinction between the ethical and unethical usage of AI.

2. Literature Review

2.1 Constructing Social Norms Through Social Media

Communication on social media among users and with online content establishes standards and norms for social integration. For students, social media platforms act as informal ecologies of information access and exchange, thus transforming the image of academic research and education processes from simple diffusion to participatory exchange of ideas and knowledge (Ross, 2017). Head et al. (2020) note that the distribution of "information and its redistribution through search and social media platforms" makes it difficult to evaluate the credibility of sources, which could, in the past, be done easily. This distribution process leads to a personalized form of information presentation, according to the traces of data left by individuals. Because people are offered different pieces of information according to their searches, and the original context is not always available, it becomes more difficult to trace the origin of information.

Consequently, social media has become a space for communicating new ideas and creating new standards for the users to adopt and implement in their everyday lives. Furthermore, the algorithmic pattern of social media caters to the individual interests of users; the algorithm functions through frequent exposure to certain content, which tends to normalize behaviors and attitudes through peer communication. Thus, social media becomes a space for diffusing norms and standards that users observe, and which they imitate, leading to internalization of these norms.

Hamed (2023) notes that social media offers unique communication features and linguistic styles, which provide users with a platform for self-expression and for sharing ideas that construct online identities. Platforms for sharing ideas appeal to the cognitive psychology of users in a way that influences them to internalize and accept the norms of social media and implement them in their own real-life cases. Among student users, it creates a sense of shared practice; witnessing their peers participating in AI prompt-sharing content reinforces a shared sense of participating in a communal practice, thereby legitimizing their intentions.

2.2 Normative Pressures and Peer-Centered Ethics in AI Use

Diffusing information and influencing the perspectives of the users by creating norms and standards via social media platforms generate a discourse relating to promoting normative peer pressure for AI use in academic practices that question the ethical constructs and behavioral responses of students. Generative AI (GenAI) tools have become significant and valuable tools for providing learning support (Ghimire et al., 2024) and for influencing the educational and instructional practices of academic institutions such as universities (Eager & Brunton, 2023). However, the popularity of GenAI tools has also increased the active use of GenAI by students to improve their academic productivity in unethical ways, which raises concerns about the ethical position of AI.

The academic community has raised concerns about the use of these GenAI tools, particularly regarding academic integrity and the potential of the use of these tools to interfere with independent research (Cotton et al., 2024; Ghimire et al., 2024; Perkins, 2023). It has been found that students have become increasingly dependent on AI for academic purposes and are using it for communication, learning, and completing academic tasks (Ursavaş et al., 2025). Lasser and Poehhacker (2025) refer to social media as the digital town square, which provides a virtual space for users to engage in extensive communication and share information. They report that, in 2020, social media was being used by 57% of European Union citizens between the ages of 16 and 74 years; currently, young adults aged 16–24 years comprise 87% of social media users.

This data validates the intention to investigate how social media promotes communication through AI use and how it exacerbates peer pressure and a fear of missing out on the usefulness of GenAI to increase productivity. This understanding from the data confirms the diffusion of information about technological innovation in the domain of academia while masking questions

about its ethical use by students through the promise of increased productivity and the provision of academic support.

2.3 Theoretical Framework: Diffusion of Innovation Theory

The study applied the diffusion of innovation (DOI) theory to investigate how social media communication and AI prompt sharing of posts for academic reasons by Chinese students influenced their behavior and normalized the unethical use of AI for academic purposes. According to Rogers (2003), diffusion is the process by which innovation is communicated through specific channels among the members of a social system over time. Accordingly, this study used DOI to analyze whether social media communication involving the sharing of AI prompts among students encourage their unethical use of AI for academic purposes. Accordingly, the paper analyses how such exposure may unfold for students by studying the DOI about AI prompts for academic use and the potential effects this DOI may have on students' motivation to use AI unethically. The paper explains the role of social media communications in promoting and normalizing the unethical use of AI.

To that end, the study aligns with the theoretical framework of the DOI theory, which posits that such practices operate through the creation and communication of information among participants to achieve mutual benefits for learning, understanding, and interpreting AI use for academic purposes (García-Avilés, 2020). The paper examines how social media communication influences students' behavior and ethical attitudes regarding the use of AI for academic purposes. Farajnezhad et al. (2021) reports that DOI is often used to investigate factors that influence the decisions of individuals relating to their adoption of technology innovations. This study used the DOI approach to evaluate how exposure to AI prompt-sharing content influenced the attitudes of the students toward the use of AI for academic purposes.

2.4 Theoretical Framework: Social Cognitive Theory

The second framework that this study employed is social cognitive theory, which defines the learning patterns of individuals that arise from the external influence of their social interactions. The theory was developed by Bandura (2001), who presents personal agency as the determining factor of individual behavior, which operates within a broad matrix of sociostructural influences that make the individual both a producer and a product of the system. The theory functions through three modes: direct individual agency, reliance on others to secure desired outcomes, and collective agency resulting from coordinating and interdependent social efforts. The theory studies human behavior as a learning response to people's social environments, while postulating a reciprocal interaction among behavioral, personal, social, and environmental factors (Schunk & Usher, 2019).

This theory, together with DOI, was used to study how students behave within the online space and in the immediate communication circle while they promote content related to AI prompt sharing. While DOI studies the sharing of new information through the technological space, students' desire to engage in such activities and their intention to engage in misconduct can be verified with social

cognitive theory. While content on social media platforms is diffused as technological innovations that influence the perception of the students, the presence of social cognition determines the intention of students to engage with these practices and the changing norms of conduct that are already established through the peer normalization of AI use. Accordingly, students' perceptions of changing norms and their intention to use AI influence their understanding of accepted behavior in an educational context and redefine their understanding of what acceptable conduct in academia. Integrating social cognitive theory and DOI provides a foundation for understanding how the personal interest and the agency of the students to engage in unethical use of AI in for academic purposes interact with the social influence of the online communication that diffuses ideas to do so. The integrated framework clarifies our understanding of student behavior in the academic space.

While DOI theory explains how AI prompt-sharing content spreads through social media and shapes students' perceptions of acceptable academic practices, social cognitive theory explains how these socially transmitted norms are internalized through moral disengagement and translated into behavioral intentions. Therefore, DOI accounts for the diffusion of AI-related behaviors across student networks, whereas social cognitive theory explains the psychological processes through which exposure becomes associated with intentions to engage in AI-assisted academic misconduct. Social cognitive theory addresses the core problems of this research, which were to investigate whether AI prompt-sharing content on social media influences how students perceive academic norms, and whether such content creates moral disengagement, which appeases students' intentions to engage in misconduct.

3. Methodology

3.1 Participant Sampling

The sampling and selection of participants for this research were conducted conveniently from undergraduate and postgraduate students at Chinese higher educational institutions. Convenience sampling was chosen for this research because of the convenience of approaching participants, the ease of data collection, and students' availability and willingness to participate. Golzar et al. (2022) identify convenience sampling as a non-probability sampling technique for choosing research participants from the target demographic because of accessibility.

In total 800 prospective participants were contacted through on-campus distribution of flyers at Chinese universities, which resulted in a sample size of 600 participants. A survey link was provided. A sample size of 600 participants exceeds the recommended standard for regression, mediation, and moderation, and improves the generalizability of the findings. The difference between the number of students approaching and those who actually responded to the survey is not large; it was observed that most participants who used the survey link responded to the questionnaire.

The inclusion and exclusion criteria for the sample were that the participants had to be enrolled students of the universities (Figure 2), had to have experience of using AI prompts for academic purposes, and had to have been exposed to social media posts that shared AI prompts. Responders who did not complete the questionnaire and those who were not enrolled in any university at the time were excluded from the research – these were only a couple of people. These inclusion and exclusion criteria were adopted to ensure that participants had direct experience relevant to the study variables and that they were already familiar with AI prompt-sharing content. Consequently, the findings should be interpreted as reflecting this specific subgroup rather than the broader population of Chinese university students. The sample size provided sufficient statistical power for the planned regression, mediation, and moderation analyses, although the use of convenience sampling limits the generalizability of the findings.

3.2 Participant Characteristics

The sampling of participants at universities resulted in a group of students ranging in age from 18 to 30+ years, with a fairly proportional gender ratio and representing various academic domains (Figure 3). Most students in the final sample of 600 were aged between 18 and 23 years; the second-largest age group was 24–26 years, with fewer participants older than 27 years.

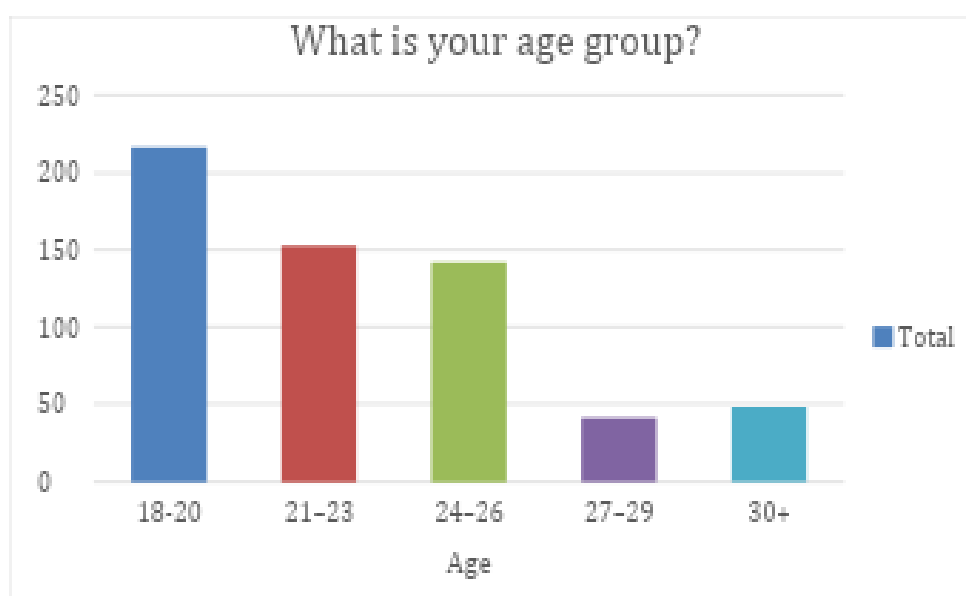


Figure 2: Age profile of the participants

Most participants were aged 18–20 (35.8%) and 21–23 (25.5%) years, together accounting for 61.3% of the sample. Participants aged 24–26 represented 23.7%, whereas only 6.8% were aged 27–29, and 8.2% were aged 30 years or older, which indicates that the sample primarily comprised young university students.

3.3 Measures and Instruments

The primary instrument used to collect participant responses is a closed and structured questionnaire-based survey that operationalized the three hypotheses of this research, namely, the exposure of students to content on AI prompt

sharing, moral disengagement, and perceived academic norms of students using AI. These hypotheses were tested to verify how the students justified an increasing intention to engage in misconduct. A questionnaire was developed for this study because no validated scale that could be used to measure exposure to content related to sharing AI prompts on social media and its connection with academic misconduct exists. In creating the items of the questionnaire, theoretical concepts of the DOI and the social cognitive theory were integrated with the literature on academic misconduct and moral disengagement.

The survey used questions with a Likert-scale format. Students' exposure to social media content on AI prompt sharing was captured using focused questions that assessed the type and frequency of engagement with posts and videos on AI prompt sharing. Psychometric checks, using the values of Cronbach's alpha, and global suitability testing were used to confirm internal consistency and the factorability of the instrument before it was administered. "Alpha was developed by Lee Cronbach in 1951 to provide a measure of the internal consistency of a test or scale; it is expressed as a number between 0 and 1" (Tavakol & Dennick, 2011, p. 53). To ensure content validity, items in the survey were designed in such a way that there was a one-to-one correspondence between the constructs and hypotheses.

The survey had four parts. Part I obtained participants' demographic information, specifically their ages and degree levels. Part II determined exposure to AI prompt-sharing content with three questions concerning: (a) How often participants viewed AI prompt-sharing tutorials/posts; (b) Participants' interaction with these types of posts (e.g., liking, saving, commenting on, or sharing them); and (c) The extent to which social media algorithms recommended AI prompt-sharing content. Part III assessed moral disengagement with three questions about participants' justification for using AI for academic tasks. Part IV investigated perceived academic norms with two questions about descriptive and injunctive peer norms regarding the use of AI.

Lastly, in Part V, participants' intentions pertaining to misconduct using AI were assessed, instead of assessing their inclination to use AI as a whole. In particular, two questions gauged their likelihood of using AI for coursework that carried grades and during time-constrained situations. Substantive questions used a five-point Likert scale, with 1 = Strongly disagree/Never/Very unlikely, depending on the item.

3.4 Data Analysis

The data analysis involved a factor-based analysis using SPSS to report the use of Kaiser-Meyer-Olkin ($KMO = .804$) and Bartlett's test of sphericity ($\chi^2(36) = 5661.11, p < .001$), as reported in the research findings. Hypothesis testing was conducted using regression-based techniques. Simple linear regression evaluated H_1 (exposure $\rightarrow$ intentions); PROCESS-style mediation analysis with bootstrapping (5,000 samples) was used to test H_2 , and moderation analysis with interaction terms was used to evaluate H_3 . Together, these approaches appropriately tested direct, indirect, and conditional roles in cross-sectional data.

Direct and indirect associations are reported using mediation analysis and also used the bootstrapped confidence intervals that established the significance of the indirect approach through moral disengagement of the participants. Furthermore, moderation analysis helped to identify the conditional relations between the perceived academic norms of students at low, medium, and high values. This approach aided in reporting simple slopes and confidence intervals that show how exposure of students to AI prompt-sharing posts leads to increasing moral disengagement as a result of strengthened permissive norms. Finally, the relational size metrics (β , R^2) and statistics (t , p) were also used to interpret the results.

3.5 Ethical Considerations

All the students who participated in this study did so voluntarily by giving informed consent electronically before completing the questionnaire. The subjects were made aware of the objectives of the study, and their confidentiality and anonymity were guaranteed; they were also informed that they could leave the survey process at any stage of participation without fear of repercussions. The personal identification details of the participants were not recorded, to maintain the privacy of the participants, and the collected data that were relevant to the research purposes were effectively removed and not retained after the research was completed.

4. Findings

The results of statistical studies, which include descriptive statistics, reliability testing, regression, mediation, and moderation analyses for hypothesis testing, are presented in this section. The research findings are based on the responses received from the 600 participants. Undergraduate students constituted the largest proportion of the sample (61%), followed by postgraduate students (24%). Participants enrolled in other programs accounted for 8%, while PhD students represented the smallest group (6%) (Figure 3).

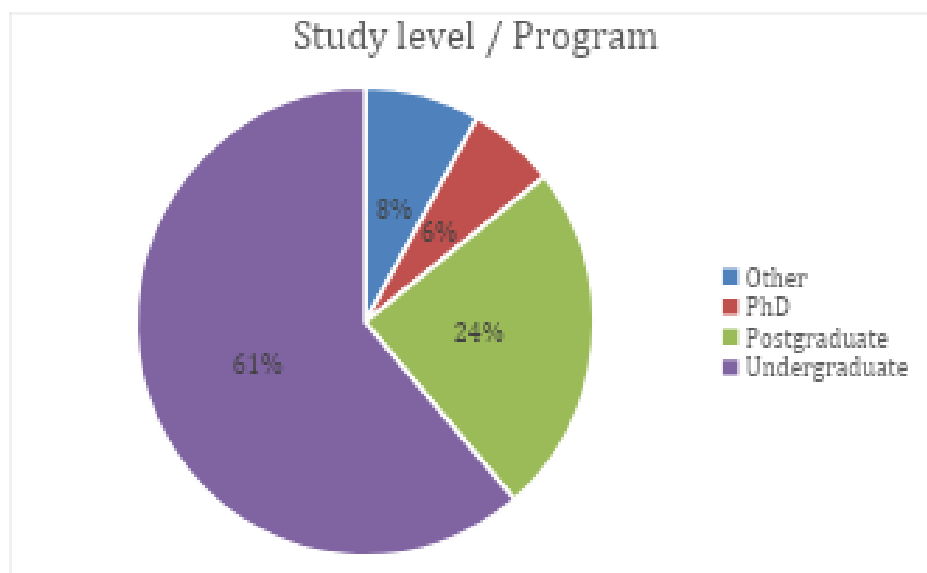


Figure 3: Levels of educational programs of participants

4.1 Reliability and Validity

Table 1: Reliability results

Scale	No. of items	Cronbach's α
Exposure to AI prompt-sharing	2	.887
Moral disengagement	3	.901
Perceived academic norms	2	.927
Intentions of academic misconduct	2	.853

Table 1 presents the reliability results for all study constructs. Cronbach's alpha values ranged from .853 to .927, indicating good to excellent internal consistency across all measurement scales. These results demonstrate that the questionnaire items reliably measured their respective constructs and were therefore suitable for subsequent statistical analyses.

Table 2: Validity results

Test	Value
Kaiser-Meyer-Olkin measure of sampling adequacy	.804
Bartlett's test of sphericity: $\chi^2(36)$	5661.11
p -value	< .001

The KMO measure of sampling adequacy is .804, which is above the recommended minimum value of .60 (Kaiser, 1974), indicating suitability for further analysis. Bartlett's test of sphericity was significant, $\chi^2(36) = 5661.11$, $p < .001$, indicating that the correlation matrix is not an identity matrix.

4.2 Hypothesis Testing

Table 3: Simple linear regression predicting intentions of academic misconduct a exposure to AI prompt-sharing content

Predictor	B	SE B	β	t	p
Constant	0.966	0.089	—	10.899	< .001
Exposure to AI prompt-sharing content	0.748	0.023	0.800	32.634	< .001
Model Summary:					
R = .800, $R^2 = .640$, Adjusted $R^2 = .640$, SE = .60187					
ANOVA: $F(1, 598) = 1064.987$, $p < .001$					

Notes. B = unstandardized regression coefficient; SE B = standard error of the unstandardized regression coefficient; β = standardized regression coefficient; R = multiple correlation coefficient; R^2 = coefficient of determination; Adjusted R^2 = adjusted coefficient of determination; F = F-statistic.

The results of the simple linear regression for the effect of exposure to AI prompt-sharing content on the intention to engage in academic misconduct by students are presented in Table 3. Overall, the regression model was found to be statistically significant, $F(1, 598) = 1064.99$, $p < .001$. The independent variable explained about 64% of the variance in the dependent variable (academic misconduct intention), $R^2 = .640$. There is a statistically significant and positive association between exposure to AI prompt-sharing content and intentions to engage in academic misconduct ($B = 0.748$, $SE = 0.023$, $\beta = 0.800$, $t = 32.63$, $p < .001$).

This result suggests that high levels of exposure to AI prompt-sharing content are significantly related to intentions to engage in academic misconduct, thus supporting H₁.

Table 4: Mediation analysis of the effect of moral disengagement on the relationship between exposure to AI prompt-sharing content and intentions to engage in academic misconduct

Predictor	B	SE B	<i>t</i>	<i>p</i>	LLCI	ULCI
Constant	1.3662	0.0955	14.311	< .001	1.1787	1.5537
Exposure to AI	0.6504	0.0247	26.324	< .001	0.6018	0.6989
Predictor	B	SE B	<i>t</i>	<i>p</i>	LLCI	ULCI
Constant	-0.0842	0.0576	-1.461	.145	-0.1974	0.0290
Exposure to AI (Direct relation, c')	0.2486	0.0189	13.143	< .001	0.2115	0.2857
Moral disengagement (b path)	0.7686	0.0213	36.071	< .001	0.7268	0.8105
Mediator	Relationship	Boot SE	BootLLCI	BootULCI		
Moral disengagement	0.4999	0.0375	0.4240	0.5716		

Notes. LLCI = lower limit of the 95% confidence interval; ULCI = upper limit of the 95% confidence interval; Boot SE = bootstrap standard error; BootLLCI = bootstrap lower limit of the 95% confidence interval; BootULCI = bootstrap upper limit of the 95% confidence interval.

According to Table 4, the findings of the mediation analysis indicate that moral disengagement was strongly predicted by exposure ($B = 0.6504$, $SE = 0.0247$, $t = 26.32$, $p < .001$), suggesting that greater exposure is linked to higher moral disengagement levels. In the second stage of the model, both exposure and moral disengagement were significant predictors of academic misconduct intentions. Misconduct intentions were strongly positively impacted by moral disengagement ($B = 0.7686$, $SE = 0.0213$, $t = 36.07$, $p < .001$), whereas the direct relationship between exposure is still significant, though weaker ($B = 0.2486$, $SE = 0.0189$, $t = 13.14$, $p < .001$). The indirect relationship between exposure and misbehavior intentions through moral disengagement is statistically significant, according to a Bootstrap analysis with 5,000 samples (relation = 0.4999, BootSE = 0.0375, 95% CI [0.4240, 0.5716], with the confidence interval not containing zero). Therefore, H₂ is supported.

Table 5: Moderation analysis of the effect of perceived academic norms on the relationship between exposure to AI prompt-sharing content and moral disengagement

Predictor	B	SE B	<i>t</i>	<i>p</i>	LLCI	ULCI
Constant	2.5513	0.1795	14.214	< .001	2.1988	2.9038
Exposure to AI	0.3055	0.0503	6.072	< .001	0.2067	0.4043
Perceived academic norms	-0.5967	0.0875	-6.817	< .001	-0.7686	-0.4248
Interaction (ExpAI × PerAcN)	0.1568	0.0201	7.817	< .001	0.1174	0.1962
<i>c</i>	R ²	F	df1	df2	<i>p</i>	
.7627	.5817	275.75	3	595	< .001	
Perceived norms (W)	Relation between ExpAI	SE	<i>t</i>	<i>p</i>	LLCI	ULCI
Low (2.00)	0.6192	0.0323	19.19	< .001	0.5558	0.6825
Moderate (4.00)	0.9328	0.0526	17.72	< .001	0.8295	1.0362
High (4.50)	1.0112	0.0609	16.61	< .001	0.8917	1.1308

Notes. *df* = degrees of freedom; *ExpAI* = exposure to AI prompt-sharing content; *PerAcN* = perceived academic norms; *W* = moderator variable.

The effect of the moderator, namely the perception of academic norms, was tested using a moderation analysis, which was performed to determine if the academic norms moderated the impact of exposure to AI prompt sharing on moral disengagement. As Table 5 shows, the total model is significant statistically, $F(3, 595) = 275.75$, $p < .001$, which explains 58.2% of the variance in moral disengagement. Exposure to AI prompt sharing has a positive impact on moral disengagement ($B = 0.3055$, $SE = 0.0503$, $t = 6.07$, $p < .001$), while the negative main effect of academic norms is significant ($B = -0.5967$, $SE = 0.0875$, $t = -6.82$, $p < .001$), meaning that stronger perceptions of academic norms were associated, generally, with lower levels of moral disengagement.

Nevertheless, the statistically significant positive interaction effect ($B = 0.1568$, $SE = 0.0201$, $t = 7.82$, $p < .001$) indicates that the positive relation between exposure and moral disengagement grew stronger with an increase in academic norms. It means that, while stronger academic norms generally mean lower moral disengagement, for those students who are more exposed to AI prompt-sharing, permissive academic norms strengthen the link between exposure and moral disengagement. The conditional effects continued to be statistically significant at low, moderate, and high levels of the moderator, as the confidence intervals did not contain zero. Thus, the results correspond to Hypothesis 3 (H_3).

4.3 Hypothesis Summary

Hypothesis	Statement	Decision
H ₁	Exposure to AI prompt-sharing content via social media is positively associated with students' intentions to engage in AI-assisted academic misconduct	Supported
H ₂	Moral disengagement mediates the relationship between exposure to AI prompt-sharing content and intentions of academic misconduct	Supported
H ₃	Perceived academic norms moderate the relationship between exposure to AI prompt-sharing content and moral disengagement, such that the relationship is stronger when norms are more permissive	Supported

5. Discussion

The purpose of this study was to examine how social media communication diffused AI-use norms in academic contexts, which was done by analyzing the relationship between exposure to AI prompt-sharing content, perceived academic norms, moral disengagement, and intentions to engage in academic misconduct. The analysis of the data of this study was guided by the framework of the DOI theory, to determine how the technology-based extended communication of students on Chinese social media platforms diffuses information on AI prompt sharing for academic reasons among students.

These findings should be considered against the background of the specific context of Chinese higher education, in which new channels, such as Bilibili, Xiaohongshu, and Douyin, have become key channels for spreading information related to AI use for academic purposes. Frequent use of these sites and increasing engagement of students in the use of GenAI have created a unique environment that allows for the development of unique norms for sharing prompts and using AI in academic settings, which differs from the norms observed in other educational environments.

The first major finding demonstrates an important positive relationship between AI prompt-sharing content exposure on social media and intentions of students to engage in AI-assisted academic misconduct. The cross-sectional nature of this study does not allow for drawing any causal conclusions, but students who reported higher levels of exposure to AI prompt-sharing content had stronger intentions to engage in AI-assisted academic misconduct. The link may be related to the way social media plays a part in making AI academic actions more visible and acceptable, according to the tenets of DOI.

This finding underscores the role of social media as a normative environment in which academic shortcuts circulate as informal advice. The frequency and visibility of AI prompt-sharing posts, whether tutorials, short videos, or algorithm-driven recommendations – create a form of observational learning, and students interpret these behaviors as acceptable, common, or even academically

strategic. In this sense, social media platforms operate as peer-driven ecosystems in which the boundaries between legitimate learning support and unethical academic assistance become blurred.

The aforementioned outcome is consistent with the research purpose because it confirms that exposure may work as the diffusion process through which newly developed, morally gray academic practices become popular among students. The results of the current study are confirmed further by the outcomes of Elom et al. (2025). They state that the level of AI awareness through exposure is inversely related to ethical behavior of students and positively correlated with intentions to engage in academic misconduct. This means that greater exposure to AI prompt-sharing content posted on social media increases the probability of academic misconduct and decreases the level of ethical awareness.

Instead of considering misconduct as a set of individual acts, the findings indicate that repeated exposure to AI prompt-sharing content contributes to the normalization of AI-assisted academic misconduct by students. The popularity of such content on various social media websites, such as Xiaohongshu, Bilibili, or Douyin, makes AI hacks easy-to-follow models of behavior. The high correlation coefficient shows that students do not necessarily cheat intentionally, but because they are constantly exposed to examples of how to use AI technology to cope with academic demands. Therefore, the first hypothesis provides more information than just creating a statistical link.

A second major conclusion drawn from the findings is that moral disengagement serves as a mediator in the correlation between exposure and intentions to engage in misconduct. Indeed, this result confirms the notion that exposure itself does not trigger misconduct; rather, it functions by changing moral reasoning; that is, the actions of the students are justified psychologically; They justify their actions through moral disengagement mechanisms that devalue the ethical aspects of using the AI tool for inappropriate purposes. Three major disengagement techniques prove to be relevant in this regard. First, diffusion of responsibility takes place when the use of AI becomes acceptable when “many other students do it.” Second, euphemistic labeling turns the act of using the AI tool into one of “getting assistance” instead of “plagiarism,” thus minimizing its moral significance. Third, advantageous comparisons portray the use of AI as a relatively milder form of misconduct than cheating on an exam by copying answers from one’s classmate.

This finding directly targets the goal of finding socio-psychological paths through which exposure leads to normalization of misconduct. It is evident from the research findings that moral disengagement acts as the mediator that links passive exposure to proactive behavior intentions. There are clear theoretical implications for such an analysis because debates about academic integrity tend to focus on policies, deterrence, and technology (Sozon et al., 2026). However, the current research reveals that misconduct today—a time of increased AI use—is closely related to changes in the moral frameworks created by the digital peer environment. Thus, institutions should be aware that the ethical boundaries of

students change daily as a result of interaction with AI and academic life. The third key finding highlights that perceptions of academic norms greatly affect the exposure-moral disengagement link. Thus, the process of exposure itself does not take place independently of the perceptions of the students about what is acceptable in a particular environment. When the learning setting allows for the extensive use of AI, there is no strong psychological barrier to prevent misconduct.

These results support the second hypothesis, namely that AI prompt-sharing videos influence the perceived academic standards and dictate the moral decisions students make about AI-assisted academic activities. The correlation between exposure to AI prompt-sharing content and the perception of AI norms by students indicates that academic misconduct arising from using AI is a result of social reinforcement—effectively, peer pressure in the digital context, which has the potential to influence students to interpret their peers' AI use behavior as a benchmark for determining their own ethical position and the appropriateness of accepting norms, because social media and the digital peer engagement environment increase the visibility of other users (Hamed, 2023; Head et al., 2020).

Accordingly, these AI usage norms function as injunctive cues rather than being simply descriptive (Schunk & Usher, 2019; Ursavaş et al., 2025), which gives them the potential to legitimize misconduct. This moderating relation connects the literature on social-media-mediated norm formation (Hamed, 2023; Ross, 2017) and adoption of AI by higher education students (Ursavaş et al., 2025) by empirically demonstrating that perceived academic norms strengthen the association between exposure to AI prompt-sharing content and moral disengagement.

Instead of portraying AI-related misconduct as an ethical lapse by individual students, this research reveals that it is rooted in the collective moral perception of students who have access to AI and digital platforms, which establishes a perceived competitiveness within the digital network of students (Lasser & Poehhacker, 2025; Ross, 2017). This view challenges the traditional approaches of educational institutions that depend only on individual-level interventions, rather than investigating the deeper motivations behind these individual actions. This finding highlights the need for institutions to create intervention strategies at the mass level, to influence and determine the perceived norms of ethical AI usage and prevent misconduct.

A comprehensive approach to the research hypotheses highlights a model that diffuses unethical AI-use norms through social media peer discourse, according to which increased exposure of students to such discussions triggers their adoption of unethical AI measures because they perceive them as accessible, familiar shortcuts. The perception of the AI norms by the students amplifies such usage by transforming their exposure into a collective behavioral pattern instead of isolated observations (Hamed, 2023; Rogers, 2003). Consequently, moral disengagement creates a psychological pathway for students to justify their unethical AI use as an accepted norm, thus creating a cumulative DOI in the

online space of social media platforms. Such actions normalize the adoption of AI-related academic misconduct, which legitimizes, psychologically and socially, the actions through peer approval and a widely accepted norm among students. The integration of this DOI model expands and enhances the theoretical interpretation of the growing role and position of AI in academia and its influence on students, specifically, its encouragement of misconduct. This model functions through the interconnection of social influence, communication theory and moral psychology (Bandura, 2001; Rogers, 2003), which enables this research to interpret the psychological influence of social media on students and their decision-making process in the context of academic misconduct.

The findings of mediation and moderation analysis highlight that, in addition to regulating the use of AI technology, there is a need for measures to decrease moral disengagement by learners and enhance ethical practices in academic pursuits. The research observations relating to interpreting the practical significance of AI in academia are substantial and expose the negative implications for students and the potential of AI to expand into a ubiquitous tool that imposes new challenges and pressures on students. The evolving use of AI in academia places students under constant surveillance for ethical misconduct while simultaneously imposing peer pressure to conform to the newly established normative standards of what the ethical use of AI comprises (Cotton et al., 2024; Perkins, 2023). While this situation creates an expectation that universities establish policies to monitor the AI use of students, it also makes it imperative that universities and faculties provide pedagogical guidance for the ethical use of AI that would prevent students from engaging in misconduct.

Doing so could help promote the AI literacy of the students, which should refer to their individual cognitive judgment of the boundaries and limitations of AI use and of participating in newly interpreted academic norms that are generalized by social media rather than academia (Eager & Brunton, 2023; Ghimire et al., 2024). As a result, the policies could shape students' moral thinking and reasoning, which would train their understanding of AI and its consequences by establishing boundaries and how the students internalize such boundaries, which signals to both instructors and students that they should actively engage in evaluating the implications of AI use. These policies could promote community-based learning in the academic space, thereby targeting the understanding of students in both individual and group settings.

The core observation of this research remains the finding that students engage in academic misconduct by using AI as a result of an integrated factorizability of social media communications, changing patterns of perceived norms and behavioral responses, and their resulting moral disengagement from standardized academic conduct. Social media communication of practices that are aimed at users of student-level age stimulates active participation of these users in generating and sharing content, to stimulate peer approval and influence others, which shifts the engagement from being passive encounter. This shift influences the behavioral attitudes of students as they continue participating in such actions, essentially under the influence of peer acceptance, which reinforces

conformity to the newly accepted norms. This shift toward active participation influences students' moral engagement with the ethical values of academia negatively and thereby creates disengagement that underscores the need for institutional reforms of instructional guidance on AI use by students.

6. Conclusion

This paper aimed to understand how social media communication diffuses AI use norms in academic contexts. Rather than viewing AI-assisted academic misconduct solely as an individual ethical issue, this study demonstrates that it is embedded in a broader social communication environment, one in which exposure to AI prompt-sharing content, perceived academic norms, and moral disengagement operate together.

By integrating DOI theory and social cognitive theory, the study extends our understanding by explaining not only how AI-related practices circulate through student networks but also how these practices become cognitively justified within academic decision-making. Consequently, the findings suggest that students may reinterpret acceptable academic behavior through their perceptions of peer practices and digital communication environments. Together, these ideas combine to present a coherent explanation of the influence of peer pressure mediated by digital communication platforms of social media, and the resulting moral cognition of the distinction between ethical and unethical practices in academia.

These findings contribute to the increasing literature on GenAI in higher education by highlighting that challenges regarding the ethics of AI use should be understood not only as issues of individual behavior but also as products of digital communication environments and shared social norms. Consequently, strategies to promote academic integrity should combine institutional AI policies with initiatives that strengthen ethical awareness, responsible AI literacy, and critical engagement with AI-related content on social media. The current research provides a unified view of the interaction between the digital communication environment, social norms, and moral reasoning with regard to the intention of students to commit academic dishonesty using AI technologies, thus providing a basis for future research in the area.

7. Limitations and future scope of research

A limitation of this study is that research samples from universities in other countries apart from mainland China were not included. By not considering the academic norms of other countries, the scope of the observation is limited. A second limitation of the study arose from the adoption of convenience sampling, which led to the potential for bias in the selection of participants, which was meticulously mitigated by studying the real opinions of the participants. Additionally, the study addressed only university-level students and did not include students at the high-school level, who might also be users of AI for academic purposes. The paper presents a limited overview of academic misconduct, because social media might equally influence other students who

were not considered in this research. Future studies could address these issues further by framing this study as a foundational guideline.

8. Conflict of Interest

The authors do not reveal any conflicts of interest in conducting this research.

Acknowledgments

The authors wish to acknowledge the use of Grammarly in the writing of this paper. This tool was used to help improve the language and grammar of the paper.

8. References

- Ahmad, Z., Soroya, M. S. & Awais, A. (2025). Understanding academic integrity beliefs, motivations and behaviors amongst postgraduate students: A comparative analysis of public and private sector universities in Punjab, Pakistan. *Journal of Academic Ethics*, 24(1), Article 22. <https://doi.org/10.1007/s10805-025-09697-x>
- Alalwan, A. A., Rana, N., Dwivedi, K. Y., & Algharabat, R. (2017). Social media in marketing: A review and analysis of the existing literature. *Telematics and Informatics*, 34(7), 1177–1190. <https://doi.org/10.1016/j.tele.2017.05.008>
- Amigud, A., & Pell, D. J. (2022). Virtue, utility and improvisation: A multinational survey of academic staff solving integrity dilemmas. *Journal of Academic Ethics*, 20(3), 311–333. <https://doi.org/10.1007/s10805-021-09416-2>
- Appel, G., Grewal, L., Hadi, R., & Stephen, A. T. (2020). The future of social media in marketing. *Journal of the Academy of Marketing Science*, 48(1), 79–95, <https://doi.org/10.1007/s11747-019-00695-1>
- Bandura, A. (2001). Social cognitive theory: An agentic perspective. *Annual Review of Psychology*, 52(1), 1–26. <https://doi.org/10.1146/annurev.psych.52.1.1>
- Cotton, D. R. E., Cotton, P. A., & Shipway, J. R. (2024). Chatting and cheating: Ensuring academic integrity in the era of ChatGPT. *Innovations in Education and Teaching International*, 61(2), 228–239. <https://doi.org/10.1080/14703297.2023.2190148>
- Eager, B., & Brunton, R. (2023). Prompting higher education towards AI-augmented teaching and learning practice. *Journal of University Teaching and Learning Practice*, 20(5), 1–19. <https://doi.org/10.53761/1.20.5.02>
- Elom, C. O., Ayanwale, M. A., Ukeje, I. O., Offiah, G. A., Umoke, C. C., & Ogbonnaya, C. E. (2025). Does AI knowledge encourage cheating? Investigating student perceptions, ethical engagement, and academic integrity in the digital age. *International Journal of Learning, Teaching and Educational Research*, 24(4), 708–729. <https://doi.org/10.26803/ijlter.24.4.33>
- Farajnezhad, S., Bodaghi Khajeh Noubar, H., & Fakhimi Azar, S. (2021). The impact of diffusion of innovation models on user behavioral intention in adopting social media marketing. *International Journal of Nonlinear Analysis and Applications*, 12(2), 1611–1632. <http://dx.doi.org/10.22075/IJNAA.2021.5291>
- García-Avilés, J. A. (2020). Diffusion of innovation. In J. van den Bulck, D. R. Ewoldsen, M. L. Mares, & E. Scharrer (Eds.), *The International Encyclopedia of Media Psychology* (pp. 1–8). John Wiley & Sons.
- Ghimire, S. N., Bhattarai, U., & Baral, R. K. (2024). Implications of ChatGPT for higher education institutions: Exploring Nepali university students' perspectives. *Higher Education Research & Development*, 43(8), 1769–1783. <https://doi.org/10.1080/07294360.2024.2366323>
- Golzar, J., Noor, S., & Tajik, O. (2022). Convenience sampling. *International Journal of Education & Language Studies*, 1(2), 72–77. <https://doi.org/10.22034/ijels.2022.162981>

- Hamed, S. M. S. H. (2023). The impact of social media on linguistic practices and cultural norms. *Journal of Arts, Literature, Humanities and Social Sciences*, 95, 276–295. <https://doi.org/10.33193/jalhss.95.2023.898>
- Head, A. J., Fister, B., & MacMillan, M. (2020). *Information literacy in the age of algorithms: Student Experiences with news and information, and the need for change*. Project Information Literacy.
- Kaiser, H.F. (1974). An index of factorial simplicity. *Psychometrika*, 39(1), 31–36. <https://doi.org/10.1007/BF02291575>
- Kotler, P., Keller, K. L., & Chernev, A. (2022). *Marketing management*. Pearson.
- Lasser, J., & Poehchacker, N. (2025). Designing social media content recommendation algorithms for societal good. *Annals of the New York Academy of Sciences*, 1548(1), 20–28. <https://doi.org/10.1111/nyas.15359>
- Lee, K. (2022). *Social media marketing for small business*. BNG Books.
- Narayanan, A. (2023, March 9). *Understanding social media recommendation algorithms*. Knight First Amendment Institute at Columbia University. <https://knightcolumbia.org/content/understanding-social-media-recommendation-algorithms>
- Perkins, M. (2023). Academic integrity considerations of AI large language models in the post-pandemic era: ChatGPT and beyond. *Journal of University Teaching and Learning Practice*, 20(2), 1–24. <https://doi.org/10.53761/1.20.02.07>
- Popescu, R.-I., Sabie, O. M., & Truşcă, M. I. (2023). The contribution of artificial intelligence to stimulate the innovation of educational services and university programs in public administration. *Transylvanian Review of Administrative Sciences*, 19(70 E), 85–108. <https://doi.org/10.24193/tras.70E.5>
- Rogers, E. M. (2003). *Diffusion of innovations* (5th ed.). Free Press.
- Ross, J. (2017). Speculative method in digital education research. *Learning, Media and Technology*, 42(2), 214–229. <https://doi.org/10.1080/17439884.2016.1160927>
- Schunk, D. H., & Usher, E. L. (2019). Social cognitive theory and motivation. In R. M. Ryan (Ed.), *The Oxford Handbook of Human Motivation* (2nd ed., pp. 10–26). Oxford University Press. <https://doi.org/10.1093/oxfordhb/9780190666453.013.2>
- Sozon, M., Alkharabsheh, O. H. M., Pok, W. F., Sia, B. C., & Chowdhury, M. A. (2026). Exploring three decades of key themes and trends in academic misconduct in higher education institutions: A literature review. *Journal of Applied Research in Higher Education*, 18(3), 815–829. <https://doi.org/10.1108/JARHE-08-2024-0440>
- Tavakol, M., & Dennick, R. (2011). Making sense of Cronbach's alpha. *International Journal of Medical Education*, 2, 53–55. <https://doi.org/10.5116/ijme.4dfb.8dfd>
- Ursavaş, Ö. F., Yalçın, Y., İslamoğlu, H., Bakır-Yalçın, E., & Cukurova, M. (2025). Rethinking the importance of social norms in generative AI adoption: Investigating the acceptance and use of generative AI among higher education students. *International Journal of Educational Technology in Higher Education*, 22(1), 1–22. <https://doi.org/10.1186/s41239-025-00535-z>
- Wong, S. S. H., Lim, S. W. H., & Quinlan, K. M. (2016). Integrity in and beyond contemporary higher education: What does it mean to university students? *Frontiers in Psychology*, 7, Article 1094. <https://doi.org/10.3389/fpsyg.2016.01094>