

Comparing Student Acceptance of ChatGPT and Claude AI Through an Extended Technology Acceptance Model

Paisal Hamid Marpaung¹, Alfiansyah Halomoan Siregar¹, Hakiki Pandapotan Hasibuan^{1,*}

¹ Universitas Muhammadiyah Tapanuli Selatan, Indonesia

* Corresponding author: Hakiki Pandapotan Hasibuan, Universitas Muhammadiyah Tapanuli Selatan, Indonesia
✉ hakiki180604@gmail.com

Copyright: © 2026 by the authors

Received: 20 July 2026 | Revised: 28 July 2026 | Accepted: 20 August 2026 | Published: 26 August 2026

Abstract

The classical Technology Acceptance Model falls short in fully explaining how students embrace generative artificial intelligence, as the impact of trust on their intention to use varies by platform. This study explores the acceptance mechanisms of ChatGPT and Claude AI by extending the Technology Acceptance Model to include Trust and Prior Experience. This study involved 312 students from the Faculty of Science and Technology at Universitas Muhammadiyah Tapanuli Selatan, selected through purposive sampling. Data were gathered using a validated Likert-scale questionnaire and analyzed with Partial Least Squares Structural Equation Modeling, incorporating multi-group analysis. Both measurement models demonstrated satisfactory validity and reliability. For ChatGPT, perceived usefulness was the primary factor influencing attitude, but this did not lead to a corresponding behavioral intention. In contrast, for Claude AI, both trust and attitude significantly influenced behavioral intention. Among the eight structural paths, four showed significant differences between platforms, with the extended model providing a better explanation for Claude AI. These results suggest that the acceptance of generative AI is dependent on the platform, offering conceptual insights into the structural assumptions of the TAM and practical guidance for platform-specific AI integration policies in higher education.

Keywords: chatgpt; claude ai; extended technology acceptance model; prior experience; trust

To cite this article: Hasibuan, H. P., Marpaung, P. H., & Siregar, A. H. (2026). Comparing Student Acceptance of ChatGPT and Claude AI Through an Extended Technology Acceptance Model. *Edumatic: Jurnal Pendidikan Informatika*, 10(2), 450–459.
<https://doi.org/10.29408/edumatic.v10i2.36272>

INTRODUCTION

Generative AI chatbots have transitioned from optional tools to standard elements in students' academic routines (Dahri et al., 2024; Zhou et al., 2024). However, much of the existing research tends to categorize AI chatbots as homogeneous technology, focusing on acceptance models for individual platforms, despite notable variations in architecture, safety alignment, and interaction design among them. This approach overlooks a critical question: does the process by which students accept a chatbot vary depending on the platform? Higher education in Indonesia provides a pertinent empirical context for exploring this issue. This is not only because it serves as a data collection site, but also because Indonesian universities experience simultaneous and competitive adoption of various generative AI platforms by the



same student body, allowing for platform-level comparisons that studies focusing on a single platform cannot achieve.

In acceptance research, an empirical contradiction regarding the role of trust prompted this comparison. [Cabrera and Neville \(2026\)](#) discovered that UK students persistently engaged with generative AI even though they openly distrusted its results, valuing convenience more than epistemic assurance. This behavior challenges the typical belief that trust is essential for continued use. This implies that trust's explanatory significance might depend on the platform, a hypothesis that has not been directly tested with the same sample.

The Technology Acceptance Model (TAM), which is based on the Theory of Reasoned Action ([Fishbein & Ajzen, 1975](#)), accounts for usage intention through the concepts of Perceived Usefulness (PU) and Perceived Ease of Use (PEOU). However, it did not initially incorporate trust as a factor in acceptance. Trust was later integrated into the model through extensions designed for uncertain environments, such as online transactions ([Gefen et al., 2003](#)) and, more recently, AI systems that may produce probabilistic and potentially incorrect results ([Shahzad et al., 2024](#)). This development prompts an inquiry that the original TAM cannot address on its own: whether trust functions as an additional determinant whose impact differs across platforms. Consequently, trust is identified here as a precursor to PU rather than a direct factor influencing Behavioral Intention (BI), based on the idea that users evaluate a system's reliability before considering its usefulness. Additionally, Prior Experience (EXP) is recognized as another precursor, influencing initial expectations regarding ease of use and trust ([Taylor & Todd, 1995](#)).

Empirical applications of TAM to generative AI consistently identify perceived usefulness (PU) and perceived ease of use (PEOU) as key predictors of intention to use ([Dahri et al., 2024](#); [Lai et al., 2023](#); [Ma et al., 2024](#)). However, there is significant variation in the role of trust, which has been found to be a dominant factor in some studies ([Balaskas et al., 2025](#); [Shahzad et al., 2024](#)) but weak or insignificant in others when usefulness and ease of use are considered ([Cabrera & Neville, 2026](#)). This variation is more likely due to differences in platform, sample, and usage context, rather than just measurement errors, as these factors can influence the importance of trust compared to usefulness. Platform characteristics, such as reliability, safety alignment, and explainability, are increasingly seen as boundary conditions for evaluating chatbots ([Bai et al., 2022](#); [Roumeliotis & Tselikas, 2023](#)). Additionally, students' epistemic evaluations of AI outputs further influence their trust assessments ([Ibrahim, 2023](#); [Kleine et al., 2025](#)). If the strength of trust is contingent on these boundary conditions, studies focusing on a single platform may not be able to differentiate between a genuinely weak trust effect and one that is strong only under certain conditions.

However, direct testing of this possibility has not yet been conducted. Current research can be categorized into distinct areas: studies that specifically estimate the TAM model for ChatGPT ([Dasian & Desriyeni, 2024](#); [Ma et al., 2024](#)), research that extends TAM by incorporating trust for individual platforms such as Perplexity ([Elfirdaus et al., 2024](#)), and comparative analyses that look at multiple platforms but focus only on brand trust or general perception, without assessing whether the structural paths of an extended TAM statistically vary across platforms ([Falebita et al., 2025](#)). None of these studies maintain consistency in respondents, constructs, and instruments while changing only the platform, which is necessary to determine if the structural relationships within TAM are inherently platform-dependent rather than just descriptively different across independently sampled groups.

This gap highlights three interconnected yet separate concepts: a "platform difference," which refers to a noticeable outcome variation between two specified platforms; a "platform characteristic," which denotes an assumed attribute like reliability or safety alignment that a platform may have; and "platform as a boundary condition or moderator," which is a formal assertion that the strength of a relationship varies depending on the platform. This study focuses

solely on the first concept, employing a comparison with the same sample and instrument. Theoretical justifications such as reliability, explainability, safety alignment, uncertainty reduction, and perceived credibility have been proposed to explain why such a difference might reasonably occur between platforms with differing design philosophies. However, since these mechanisms were not directly assessed in this study, any interpretation involving them remains a plausible explanation rather than a proven causal link.

ChatGPT and Claude AI were selected because of their documented differences in design focus. ChatGPT, developed on the GPT framework, is widely acknowledged for its general-purpose functionality, although it raises issues such as enabling plagiarism (Shoufan, 2023). In contrast, Claude AI employs a Constitutional AI strategy that prioritizes safety alignment and output dependability (Bai et al., 2022). This documented distinction serves solely as a basis for anticipating varied evaluations of the platforms, rather than presuming that ChatGPT's utility will prevail or that Claude AI's trustworthiness will be more significant. The determination of which aspect is more prominent is an empirical matter resolved by structural findings, not pre-assumed. Studies based on the TAM model involving ChatGPT among Indonesian students are becoming more established (Dasian & Desriyeni, 2024; Elfirdaus et al., 2024), yet they focus on a single platform independently. Broader comparative research primarily relies on international samples (Cabrera & Neville, 2026). The acceptance process of generative AI among Science and Technology students at private, religious universities in regional Indonesia, such as Universitas Muhammadiyah Tapanuli Selatan (UMTS), has not yet been empirically investigated.

In this research, an extended version of the TAM is employed to examine how PEOU, PU, Trust, and Prior Experience collectively influence attitudes and intentions toward ChatGPT and Claude AI. The study compares acceptance patterns between these two platforms among UMTS Science and Technology students by utilizing the same sample and instrument design. Multi-group analysis is used to statistically determine which structural paths vary between platforms and remain consistent. This study does not introduce any new constructs; instead, its innovation lies in redefining trust as a platform-specific factor rather than a universal determinant of acceptance. This study provides a conceptual contribution to the boundary conditions of the TAM and offers a practical foundation for the UMTS Faculty of Science and Technology to develop AI integration policies tailored to different platforms.

METHOD

This study utilized a quantitative cross-sectional explanatory approach with PLS-SEM, as noted by Lim (2025). This study focuses on explanation and confirmation, aiming to test a predetermined, theory-based set of pathways rather than investigating an undefined framework. PLS-SEM was selected over CB-SEM for several reasons: some indicators did not meet the criteria for multivariate normality (with skewness and kurtosis beyond ± 1); the model's complexity, involving eight paths and seven constructs, aligns well with the robustness of PLS-SEM; and importantly, PLS-SEM inherently accommodates PLS MGA for comparing entire structural models across different platforms, which is something that the multi-group extension manages with less flexibility (Hair et al., 2019).

The population comprised 1,418 Faculty of Science and Technology students at UMTS; purposive sampling required prior use of both platforms. An initial $n = 312$ was estimated using Slovin's formula (5% margin of error). Since Slovin's formula is not PLS SEM specific, it was cross-checked against a power-based benchmark (largest predictor count for one endogenous construct = 3, minimum $f^2 = 0.20$, $\alpha = .05$, power = 80%), yielding a minimum requirement of 155 per model, well below the achieved 312. Each respondent evaluated both platforms, so every construct was measured twice using identically worded instruments differing only in platform name, and presentation order was randomized to reduce order effects. Because

ChatGPT and Claude AI data therefore come from the same respondents rather than independent groups, the two platform models were analyzed separately and compared using PLS MGA's permutation procedure, which does not assume between-group independence.

The study examined seven constructs within the extended TAM Model: Perceived Ease of Use (PEOU), Perceived Usefulness (PU), Attitude Toward Use (ATU), Behavioral Intention (BI), Actual Use (AU), Trust (TR), and Prior Experience (EXP). PEOU refers to students' perceptions of how easy ChatGPT or Claude AI is to use, whereas PU reflects the perceived benefits of each platform for academic activities. ATU represents students' positive or negative evaluations of platform use, BI indicates their intention to continue using it, and AU reflects their actual use in academic activities. TR refers to confidence in the reliability and credibility of the platform, whereas EXP represents previous exposure and familiarity, measured through self-reported frequency of use and perceived familiarity. All constructs were measured using a five-point Likert scale with three indicators each, except for AU, which consisted of two indicators. The same instrument was used for both platforms, differing only in the platform name.

All constructs used a 5 point Likert scale with three items each (two for AU), adapted from established TAM and trust indicators. Prior Experience was operationalized as self-reported frequency of use and perceived familiarity, not elapsed time. The instrument was reviewed by two information systems lecturers for content validity and pilot tested with 30 students outside the sample before distribution. Data were analyzed using SmartPLS 4. The measurement model was evaluated via outer loadings (≥ 0.70), AVE (≥ 0.50), composite reliability and Cronbach's alpha (≥ 0.70), and discriminant validity (HTMT ≤ 0.90). Measurement invariance (MICOM) was tested before any structural comparisons. The structural model was evaluated using bootstrapped path coefficients (5,000 subsamples, $\alpha=0.05$), R^2 , f^2 , Q^2 predict, and VIF (< 3.3). The cross-platform comparison used permutation-based PLS MGA (5,000 permutations), the non-parametric variant suited to this within-subject design. A path was deemed significantly different at $p \leq 0.05$, interpretable as genuine only because invariance was confirmed beforehand.

RESULTS AND DISCUSSION

Results

This section presents the empirical results following the analytical sequence described in the Methods section: measurement model, measurement invariance, structural model, and multigroup comparison of structural paths between the two platforms. In the initial phase, convergent validity is evaluated using outer loadings (≥ 0.70) and Average Variance Extracted (AVE) (≥ 0.50). All outer loadings for both models exceed the 0.70 benchmark.

Table 1. Summary of ave ($\geq 0,50$), composite reliability ($\geq 0,70$), and cronbach's alpha ($\geq 0,70$)

Construct	AVE ChatGPT	AVE Claude	CR ChatGPT	CR Claude	α ChatGPT	α Claude
ATU	0,750	0,653	0,900	0,748	0,846	0,737
AU	0,730	0,750	0,843	0,715	0,692	0,674
BI	0,702	0,724	0,876	0,812	0,789	0,809
EXP	0,753	0,764	0,901	0,869	0,886	0,847
PEOU	0,619	0,620	0,829	0,693	0,694	0,692
PU	0,672	0,685	0,860	0,776	0,757	0,770
TR	0,648	0,721	0,846	0,830	0,727	0,809

Table 2. Heterotrait-monotrait ratio ($htmt \leq 0.90$)

Construct Pair	HTMT ChatGPT	HTMT Claude	Construct Pair	HTMT ChatGPT	HTMT Claude
AU ↔ ATU	0,146	0,522	PU ↔ ATU	0,394	0,601
BI ↔ ATU	0,060	0,706	PU ↔ AU	0,104	0,336
BI ↔ AU	0,735	0,769	PU ↔ BI	0,094	0,425
EXP ↔ ATU	0,122	0,181	PU ↔ EXP	0,085	0,274
EXP ↔ AU	0,021	0,113	PU ↔ PEOU	0,618	0,767
EXP ↔ BI	0,058	0,119	TR ↔ ATU	0,100	0,418
PEOU ↔ ATU	0,159	0,443	TR ↔ AU	0,093	0,245
PEOU ↔ AU	0,177	0,248	TR ↔ BI	0,166	0,322
PEOU ↔ BI	0,159	0,283	TR ↔ EXP	0,030	0,239
PEOU ↔ EXP	0,130	0,660	TR ↔ PEOU	0,832	0,863
			TR ↔ PU	0,695	0,887

Table 2 verifies the discriminant validity of both models, with all HTMT ratios remaining under the 0.90 limit. However, the Claude AI model exhibits significantly elevated ratios for pairs that include trust, especially TR↔PEOU (0.863) and TR↔PU (0.887). This suggests a closer relationship between trust and cognitive evaluative constructs on this platform, while still maintaining the distinctiveness of the constructs.

Tabel 3. R-square

ChatGPT			Claude AI		
	R-square	R-square adjusted		R-square	R-square adjusted
ATU	0,101	0,095	ATU	0,219	0,213
AU	0,285	0,283	AU	0,334	0,332
BI	0,000	-0,003	BI	0,305	0,303
PEOU	0,010	0,007	PEOU	0,272	0,270
PU	0,300	0,295	PU	0,538	0,535
TR	0,000	-0,003	TR	0,039	0,036

With the invariance established, the structural model was evaluated. Table 3 shows that the Claude AI model explains substantially more variance than the ChatGPT model across most endogenous constructs, most notably Behavioral Intention ($R^2 = 0.305$ vs. 0.000) and Perceived Usefulness ($R^2 = 0.538$ vs. 0.300). Trust remains weakly explained in both models, although it is comparatively higher for Claude AI (0.039 vs. 0.000). Effect sizes (f^2) follow a similar pattern for ATU→BI and EXP→PEOU (large for Claude AI, $f^2 = 0.439$ and 0.374, versus negligible for ChatGPT, $f^2 = 0.000$ and 0.010), but PEOU→PU is small and comparable across platforms ($f^2 = 0.046$ vs. 0.036), indicating that this particular contribution does not depend on the platform. Q^2 prediction values were mostly positive and meaningful for Claude AI, whereas several ChatGPT indicators (PEOU, TR, and all BI and AU items) showed near zero or negative values, signaling weak predictive relevance. SRMR falls within acceptable limits for both models (0.112 ChatGPT; 0.132 Claude AI). Consistent with the reviewer's caution, these f^2 , Q^2 , and SRMR results are treated as supporting diagnostics rather than direct evidence that one platform's acceptance mechanism is substantively stronger.

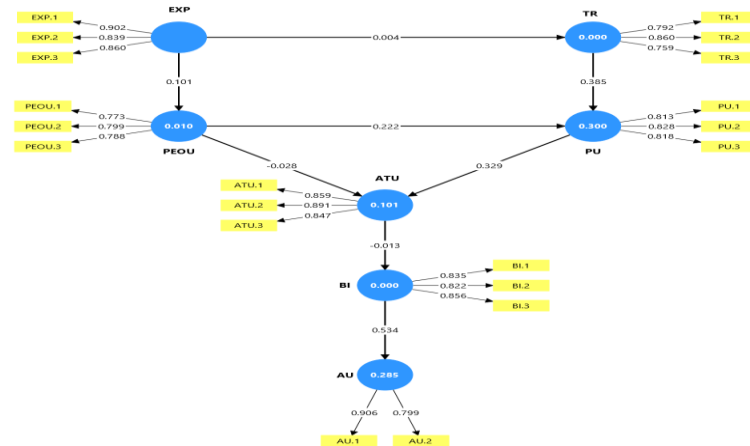


Figure 1. Multi group permutation path coefficient

Figure 1 presents the structural path model for ChatGPT, showing standardized path coefficients and R^2 values for each construct, and reveals a clear asymmetry in explanatory strength, Prior Experience has a negligible effect on Trust ($\beta = 0.004$) and only a weak effect on PEOU ($\beta = 0.101$), leaving both PEOU ($R^2 = 0.010$) and TR ($R^2 = 0.000$) almost entirely unexplained plausibly reflecting ChatGPT's status as a now-familiar tool with limited variance in prior experience across respondents whereas Perceived Usefulness is comparatively well explained ($R^2 = 0.300$), driven by both Trust ($\beta = 0.385$) and PEOU ($\beta = 0.222$), and in turn exerts a moderate influence on Attitude toward Use ($\beta = 0.329$) while PEOU's direct effect on Attitude is negligible ($\beta = -0.028$), indicating that attitude formation is driven by usefulness rather than ease of use most notably, Attitude barely explains Behavioral Intention ($\beta = -0.013$, $R^2 = 0.000$), suggesting favorable attitudes do not translate into stronger intention to use ChatGPT and that intention is likely shaped by factors outside the TAM chain such as habit or necessity, yet Behavioral Intention still strongly predicts Actual Use ($\beta = 0.534$, $R^2 = 0.285$) together showing that ChatGPT's acceptance mechanism operates unevenly, with the usefulness formation segment (PEOU/TR $\rightarrow$ PU $\rightarrow$ ATU) following classical TAM logic while the attitude intention link central to TAM is empirically inert, pointing to a usage pathway that bypasses attitudinal mediation once use becomes habitual.

Table 4. Micom

	Original correlation	Correlation permutation mean	5.0%	Permutation p value
ATU	0,997	0,998	0,995	0,161
AU	1,000	0,999	0,995	0,775
BI	1,000	0,999	0,998	0,921
EXP	0,993	0,998	0,992	0,064
PEOU	0,998	0,998	0,992	0,449
PU	1,000	0,999	0,998	0,834
TR	0,996	0,998	0,995	0,155

Following the confirmation of measurement adequacy, multi-group invariance testing was executed to ensure data equivalence across platforms. Table 4 shows that all constructs satisfy configural and compositional invariance, with permutation p values exceeding 0.05 for every construct, establishing that the two platform models are sufficiently comparable and providing the methodological basis for the cross-platform comparison reported next.

Table 5. PLS multi-group analysis (path coefficient comparison $p \leq 0.05$)

	β ChatGPT	β Claude AI	$\Delta\beta$	2.5%	97.5%	T statistik	p value
ATU > BI	-0.013	0.552	-0.566	-0.164	0.161	5.882	0.000
BI > AU	0.534	0.578	-0.044	-0.129	0.132	16.722	0.502
EXP > PEOU	0.101	0.522	-0.421	-0.178	0.164	7.533	0.000
EXP > TR	0.004	0.198	-0.194	-0.154	0.163	2.691	0.013
PEOU > ATU	-0.028	0.088	-0.116	-0.200	0.207	0.373	0.258
PEOU > PU	0.222	0.167	0.056	-0.173	0.188	4.334	0.589
PU > ATU	0.329	0.413	-0.084	-0.200	0.179	7.594	0.398
TR > PU	0.385	0.616	-0.231	-0.190	0.175	10.446	0.014

Table 5 reports the central empirical results of this study. Four of the eight structural paths differ significantly between the ChatGPT and Claude AI models: ATU→BI ($\Delta\beta = -0.566$, $p = 0.000$), EXP→PEOU ($\Delta\beta = -0.421$, $p = 0.000$), EXP→TR ($\Delta\beta = -0.194$, $p = 0.013$), and TR→PU ($\Delta\beta = -0.231$, $p = 0.014$); all four paths are stronger in the Claude AI model. The remaining four paths, BI→AU, PEOU→ATU, PEOU→PU, and PU→ATU, do not differ significantly across platforms, indicating that the usefulness attitude intention chain central to classical TAM remains comparatively stable regardless of platform, whereas the paths involving Trust and Prior Experience are platform-dependent.

Taken together, the measurement models demonstrate adequate validity and reliability for both platforms, measurement invariance is established, and the multi-group comparison provides direct statistical evidence that four of eight structural paths, primarily those involving Trust and Prior Experience, differ significantly between ChatGPT and Claude AI, while the core usefulness based pathway of TAM remains invariant across platforms.

Discussion

The central finding is that four of the eight structural paths in the extended TAM differ significantly between the ChatGPT and Claude AI models, namely, the relationships between ATU and BI, EXP and PEOU, EXP and TR, and TR and PU. In contrast, the relationships between PEOU and ATU, PEOU and PU, PU and ATU, and BI and AU remain statistically invariant. Therefore, generative AI acceptance is only partially platform-dependent. TAM's classical usefulness-driven pathway remains stable across platforms, whereas the relationships involving TR and EXP are platform-conditional. This pattern extends rather than challenges the TAM, as the original attitudinal chain linking PEOU and PU to ATU and BI is neither contradictory nor weakened. Instead, the findings challenge the implicit assumption inherited from single-platform applications that antecedents such as TR and EXP operate with uniform strength, regardless of the technology being evaluated. The extended TAM is therefore supported at its structural core, while requiring further elaboration at the antecedent level.

A plausible, though untested, explanation for the stronger trust-related relationships in Claude AI is its Constitutional AI design, which emphasizes safety alignment and output reliability (Bai et al., 2022). As perceived reliability and safety alignment were not directly measured, this interpretation should be regarded as a theoretical explanation consistent with the observed pattern, rather than a demonstrated causal mechanism. The findings support significant cross-platform differences but do not establish that a specific design attribute moderates these relationships. The non-significant relationships are equally important because their combination with the significant findings constitutes the study's main theoretical contribution. The relationships between PEOU and PU, PU and ATU, and BI and AU do not differ significantly across platforms, indicating that once usefulness and ease of use are

established, their influence on attitude and behavior operates similarly across platforms, consistent with the TAM's core logic (Fishbein & Ajzen, 1975). The relationship between PEOU and ATU warrants particular attention because PEOU shows a weak and statistically indistinguishable effect on ATU in both models ($\beta = -0.028$ for ChatGPT; $\beta = 0.088$ for Claude AI). This suggests that ATU is influenced more strongly by PU than PEOU, consistent with evidence that the direct attitudinal effect of PEOU may diminish as technology becomes more familiar (Taylor & Todd, 1995). Thus, PEOU appears to contribute primarily through its effect on PU, which remains platform-invariant. Overall, platform dependence is concentrated in antecedent relationships involving TR and EXP rather than in TAM's attitudinal core, making the combined pattern of invariance and divergence the central theoretical contribution of this study.

The near-zero R^2 values for BI and TR in the ChatGPT model (0.000 for both compared with 0.305 and 0.039 for Claude AI) represent substantive findings rather than methodological weaknesses. Because ChatGPT is already familiar and routinely used, the limited variation in EXP across respondents may have reduced its ability to explain PEOU and TR, which is consistent with evidence that the explanatory role of experience decreases as technology use becomes habitual (Taylor & Todd, 1995). This suggests that BI toward ChatGPT may also be influenced by unmeasured factors, such as habit or social norms. The stronger role of TR in Claude AI is consistent with studies emphasizing trust in chatbot adoption (Balaskas et al., 2025; Shahzad et al., 2024), whereas the stronger role of PU in ChatGPT differs from that of Falebita et al. (2025), who identified brand-level trust as a primary determinant. This difference may reflect variations in operationalization, sample characteristics, and platform maturity, rather than a direct contradiction. Furthermore, the invariant relationships among PEOU, PU, and ATU extend previous single-platform studies among Indonesian students by showing that this core TAM relationship also generalizes beyond ChatGPT (Dasian & Desriyeni, 2024; Elfirdaus et al., 2024).

The scientific contribution here is that extending the TAM to generative AI reveals that structural relationships are not uniformly invariant across platforms; the factors influencing usefulness vary by platform, whereas attitudinal outcomes do not. Since this conclusion is based on PLS-MGA rather than a formally tested moderation model, it is presented as a documented cross-platform difference in structural paths rather than as evidence that platform design orientation moderates TAM as a formal construct. In practice, AI integration efforts at UMTS's Faculty of Science and Technology could be tailored by providing platform-specific hands-on training that highlights task usefulness for ChatGPT and offering transparent demonstrations of output reliability to foster trust in Claude AI. For developers, platforms focused on safety should make trustworthiness apparent, whereas general-purpose platforms should continue to enhance versatility and responsiveness.

This study is constrained by its cross-sectional design, purposive sampling, and the context of a single institution and faculty, which limits its generalizability. The characteristics of the platform used to interpret findings related to trust, such as reliability, safety alignment, and explainability, have not been assessed directly. Future studies should aim to replicate this design in a wider range of institutions, utilize longitudinal or experimental designs to determine causal relationships, and directly evaluate platform characteristics to verify whether they formally influence the pathways identified here as dependent on the platform.

CONCLUSION

This study investigates whether the structural relationships within the extended TAM differ between ChatGPT and Claude AI among UMTS Science and Technology students. These findings indicate that these differences are selective rather than uniform. PU emerged as the primary factor influencing ChatGPT acceptance, whereas trust played a more prominent

role in Claude AI acceptance. Significant differences were identified in the relationships between ATU and BI, EXP and PEU, EXP and TR, and TR and PU. In contrast, the relationships among PEU, PU, and ATU, as well as between BI and AU, remained relatively consistent across both platforms. The main contribution of this study is evidence that the structural relationships within an extended TAM are not entirely stable across generative AI platforms, offering a more nuanced understanding than simply concluding that each platform has different acceptance factors. However, because platform design characteristics were not directly measured, this study does not claim that Constitutional AI or any specific design feature caused the stronger role of Trust in the Claude AI model. The conclusions are also limited by the cross-sectional design, purposive sampling, and single-institution context. Future research should, therefore, examine the model across broader university populations, employ longitudinal or experimental designs, and directly measure platform-specific characteristics such as perceived reliability, safety, and user experience.

REFERENCES

- Bai, Y., Kadavath, S., Kundu, S., Askill, A., Kernion, J., Jones, A., Chen, A., Goldie, A., Mirhoseini, A., & McKinnon, C. (2022). Constitutional AI: Harmlessness from AI feedback. *arXiv preprint*. <https://doi.org/10.48550/arXiv.2212.08073>
- Balaskas, S., Tsiantos, V., Chatzifotou, S., & Rigou, M. (2025). Determinants of ChatGPT adoption intention in higher education: Expanding on TAM with the mediating roles of trust and risk. *Information*, 16(2), 82. <https://doi.org/10.3390/info16020082>
- Cabrera, C., & Neville, R. (2026). Widely used but barely trusted: Understanding student perceptions on the use of generative AI in higher education. *Studies in Higher Education*. <https://doi.org/10.1080/13603108.2025.2595453>
- Dahri, N. A., Yahaya, N., Al-Rahmi, W. M., Aldraiweesh, A., Alturki, U., Almutairy, S., & Soomro, R. B. (2024). Extended TAM based acceptance of AI-powered ChatGPT for supporting metacognitive self-regulated learning in education: A mixed-methods study. *Heliyon*, 10(8), e29317. <https://doi.org/10.1016/j.heliyon.2024.e29317>
- Dasian, R. S. I., & Desriyeni, D. (2024). Acceptance of ChatGPT technology among students: A descriptive study of the TAM model on students of the Informatics Engineering Study Program, Padang State University. *Student Research Journal*, 2(2), 178-201. <https://doi.org/10.55606/jsr.v2i2.2847>
- Elfirdaus, I., Suryanto, T. L. M., & Pratama, A. (2024). Evaluation of student acceptance of the Perplexity application as a learning aid using a simplified Technology Acceptance Model (TAM). *Journal of Informatics and Applied Electrical Engineering*, 12(3). <https://doi.org/10.23960/jitet.v12i3.4803>
- Falebita, O. S., Abah, J. A., Asanre, A. A., Abiodun, T. O., Ayanwale, M. A., & Ayanwoye, O. K. (2025). Determinants of chatbot brand trust in the adoption of generative artificial intelligence in higher education. *Education Sciences*, 15(10), 1389. <https://doi.org/10.3390/educsci15101389>
- Fishbein, M., & Ajzen, I. (1975). *Belief, attitude, intention, and behavior: An introduction to theory and research*. Addison-Wesley.
- Gefen, D., Karahanna, E., & Straub, D. W. (2003). Trust and TAM in online shopping: An integrated model. *MIS Quarterly*, 27(1), 51-90. <https://doi.org/10.2307/30036519>
- 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>
- Ibrahim, H. (2023). Perception, performance, and detectability of conversational artificial intelligence across 32 university courses. *Scientific Reports*, 13(1). <https://doi.org/10.1038/s41598023389643>

- Kleine, A.K., Schaffernak, I., & Lermer, E. (2025). Exploring predictors of AI chatbot usage intensity among students: Within and between person relationships based on the Technology Acceptance Model. *Computers in Human Behavior: Artificial Humans*, 3, 100113. <https://doi.org/10.1016/j.chbah.2024.100113>
- Lai, C. Y., Cheung, K. Y., & Chan, C. S. (2023). Exploring the role of intrinsic motivation factors in ChatGPT adoption to support active learning: An extension of the technology acceptance model. *Computers and Education: Artificial Intelligence*, 5, 100178. <https://doi.org/10.1016/j.caeai.2023.100178>
- Lim, W. M. (2025). What is quantitative research? An overview and guidelines. *Australasian Marketing Journal*, 33(3), 325-348. <https://doi.org/10.1177/14413582241264622>
- Ma, J., Wang, P., Li, B., Wang, T., Pang, X. S., & Wang, D. (2024). Exploring user adoption of ChatGPT: A technology acceptance model perspective. *International Journal of Human-Computer Interaction*. <https://doi.org/10.1080/10447318.2024.2314358>
- Roumeliotis, K. I., & Tselikas, N. D. (2023). ChatGPT and Open-AI models: A preliminary review. *Future Internet*, 15(6), 192. <https://doi.org/10.3390/fi15060192>
- Shahzad, M. F., Xu, S., & Javed, I. (2024). ChatGPT awareness, acceptance, and adoption in higher education: The role of trust as a cornerstone. *International Journal of Educational Technology in Higher Education*, 21(1), Article 46. <https://doi.org/10.1186/s41239-024-00478-x>
- Shoufan, A. (2023). Exploring students' perceptions of ChatGPT: Thematic analysis and follow-up survey. *IEEE Access*, 11, 38805-38818. <https://doi.org/10.1109/ACCESS.2023.3268224>
- Taylor, S., & Todd, P. A. (1995). Understanding information technology usage: A test of competing models. *Information Systems Research*, 6(2), 144-176. <https://doi.org/10.1287/isre.6.2.144>
- Zhou, X., Teng, D., & Al-Samarraie, H. (2024). The mediating role of generative AI self-regulation on students' critical thinking and problem solving. *Education Sciences*, 14(12), 1302. <https://doi.org/10.3390/educsci14121302>