

ETHICAL ADOPTION OF GENERATIVE AI IN HIGHER EDUCATION: A PRELIMINARY FATE – UTAUT STUDY OF BEHAVIOURAL INTENTION AND ETHICAL USAGE

Ravikiran N.R.

Associate Professor & Head
Department of Commerce, Government First Grade College, Cox Town, Frazer Town Post,
Bangalore, India.

Hemanth Kumar S.

Professor
Area Chair & IQAC Head, Faculty of Management Studies, CMS Business School, JAIN
(Deemed –to-be University), Bangalore, India.

Ravikumar R.

Associate Professor
Department of Commerce, Government First Grade College, Harohalli, India.

ABSTRACT

Generative artificial intelligence (GenAI) has rapidly entered higher education, but its educational legitimacy now depends not only on usefulness and ease of use, but also on whether students perceive such systems as fair, accountable, transparent, and ethically governable. This study revises and strengthens an earlier draft by developing a theoretically aligned, transparently reported preliminary empirical manuscript that integrates the Fairness, Accountability, Transparency, and Ethics (FATE) lens with the Unified Theory of Acceptance and Use of Technology (UTAUT) to examine students' behavioural intention to use GenAI and their self-reported ethical usage behaviour. The study uses a simulated dataset of 450 higher education students and 33 Likert-type indicators spanning FATE, UTAUT, behavioural intention, ethical usage behaviour, and AI literacy. Because the available dataset is simulated and not field collected, the study is explicitly positioned as a preliminary model-development and measurement-validation exercise rather than as a definitive causal test. Reliability analysis showed acceptable internal consistency across all constructs (Cronbach's alpha values above .80). Descriptive results indicated high perceived usefulness, ease of use, and facilitating conditions, but only moderate evaluations of accountability and transparency. However, structural tests revealed weak cross-construct relationships, with the multivariable behavioural intention model explaining little variance and failing to produce stable statistically significant paths. Supplementary analyses for demographic differences, moderation, and robustness similarly yielded limited evidence of structural effects. These findings are substantively important because they reveal an ethical–adoption measurement gap: students may endorse GenAI-related items within constructs, yet the broader nomological network remains underdeveloped when theoretical paths are not embedded into the data-generating structure. The study therefore contributes a more honest manuscript architecture, a refined set of hypotheses, a defensible methodological narrative, and a policy-oriented discussion of why future higher education research must combine stronger measurement design, authentic field data, and institution-specific AI governance. The paper concludes with implications for universities, researchers, and conference reviewers evaluating early-stage GenAI adoption studies.

Keywords: generative artificial intelligence, higher education, FATE, UTAUT, AI literacy, behavioural intention, ethical usage behaviour

1. INTRODUCTION

Generative artificial intelligence has moved from novelty to routine presence in higher education within an unusually short time. Students now use large language models and related GenAI systems for idea generation, summarisation, translation, coding support, feedback, drafting, and study planning. Universities, meanwhile, are attempting to decide whether these tools should be discouraged, normalised, regulated, or integrated into pedagogy in a more structured way. The emerging literature shows that higher education institutions are no longer dealing with a purely technological adoption problem; they are confronting a governance problem, a pedagogical problem, and an ethical problem at the same time. Recent reviews argue that GenAI in higher education should be examined through both its affordances and its risks, especially with regard to academic integrity, privacy, explainability, bias, and human oversight (Bond et al., 2024; Munaye et al., 2025; Wang et al., 2024; Zhuang et al., 2025).

Early technology-adoption research often relied on functional predictors such as perceived usefulness, perceived ease of use, performance expectancy, and social influence. Those constructs remain important. Indeed, studies of student interaction with ChatGPT continue to show that perceived usefulness, effort expectancy, trust, and positive attitudes influence intention to use AI systems in academic work (Shuhaiber et al., 2025; Rizun et al., 2026). Yet a purely functional explanation is increasingly incomplete. Students may perceive a system as useful while simultaneously doubting whether it is fair, whether its outputs can be trusted, whether its decision logic is understandable, whether sensitive data are protected, and whether using it complies with academic norms. In other words, adoption and ethical legitimacy are related but not identical constructs.

This tension is why the FATE framework has become relevant to educational research. Memarian and Doleck (2023) show that higher education research on AI ethics has been unevenly distributed across fairness, accountability, transparency, and ethics, with fairness receiving more attention than the other dimensions and with an overall need for more longitudinal, reproducible, and empirically grounded work. Parallel evidence from policy-oriented studies indicates that universities are adopting “open but cautious” positions toward GenAI, emphasising ethical usage, accuracy, data privacy, discipline-specific policy, and staff/student support resources (Wang et al., 2024). The implication is clear: students’ behavioural intention to use GenAI cannot be fully understood without considering the ethical conditions under which adoption becomes educationally acceptable.

A second strand of recent literature concerns AI literacy. Studies increasingly suggest that AI literacy is not merely a technical competence; it includes critical evaluation, awareness of limitations, understanding of bias, and the ability to verify and appropriately use AI-generated outputs (Chen et al., 2024; van der Linde et al., 2025). Higher AI literacy appears to shape how students interact with GenAI systems in writing and learning tasks, often distinguishing more reflective and collaborative usage patterns from more passive or superficial use (Kim et al., 2025). This makes AI literacy a theoretically meaningful moderator in any model seeking to connect ethical perceptions with actual responsible use.

The present paper addresses these issues through a revised and more rigorous version of an earlier draft. The original draft proposed an integrated FATE–UTAUT model but suffered from major alignment problems between theory, measurement, and analysis. It overstated

structural findings, mentioned constructs that were not actually tested, and used a dataset whose cross-construct correlations were too weak to support the claims being made. Rather than masking these weaknesses, the present manuscript adopts a transparent corrective strategy. It reframes the study as a preliminary empirical exercise using a simulated dataset to examine the feasibility of the proposed measurement architecture, the descriptive profile of key constructs, and the extent to which the available data can sustain a plausible nomological model.

The revised paper makes four contributions. First, it consolidates FATE and UTAUT within a more coherent higher-education GenAI framework focused on behavioural intention and ethical usage behaviour. Second, it presents a cleaner hypothesis structure that distinguishes what is theoretically expected from what is empirically supported. Third, it offers an honest methodological account of simulated data use, avoiding the common error of presenting simulation-based findings as field-validated causal evidence. Fourth, it derives policy and research implications from both the strengths and the weaknesses of the current evidence base.

Accordingly, the study is guided by the following research questions: (1) How do students in the present dataset score on key FATE and UTAUT dimensions related to GenAI use in higher education? (2) Do FATE and UTAUT constructs meaningfully explain behavioural intention to use GenAI? (3) Does behavioural intention translate into ethical usage behaviour? and (4) Does AI literacy strengthen the relationship between ethical perceptions and behavioural intention or ethical usage behaviour? By answering these questions cautiously and transparently, the paper aims to provide a conference-ready manuscript that is methodologically credible, theoretically grounded, and substantially improved over the earlier version.

2. LITERATURE REVIEW AND HYPOTHESIS DEVELOPMENT

2.1 FATE and the ethical turn in higher education AI research

The spread of GenAI has intensified demands for institutionally credible ethical frameworks. FATE provides one of the most widely used umbrellas for discussing AI ethics because it captures four concerns central to educational use. Fairness relates to equitable treatment and the avoidance of discriminatory outputs. Accountability concerns responsibility for harms, errors, and misuse. Transparency addresses explainability and intelligibility of system outputs and data practices. Ethics, more broadly, speaks to integrity, privacy, moral acceptability, and alignment with human values. In their systematic review, Memarian and Doleck (2023) conclude that higher education scholarship still lacks balanced empirical attention across all four dimensions and often treats these principles more as normative aspirations than as measured student perceptions.

For higher education, this matters because student adoption of GenAI is situated within assessment regimes, institutional codes, and unequal digital competencies. Policy analysis of top universities suggests that institutions increasingly frame GenAI as something that must be governed through transparent guidance, explicit expectations, and context-sensitive evaluation strategies rather than through blanket approval or blanket prohibition (Wang et al., 2024). Systematic reviews similarly show that data privacy, opacity, bias, and accountability are among the most persistent ethical concerns when GenAI tools are used in educational settings (Munaye et al., 2025; Yan et al., 2025). On this basis, students who perceive GenAI systems to be fairer, more accountable, more transparent, and more ethically aligned should, in theory, display stronger behavioural intention to use them responsibly.

H1: Perceived fairness is positively associated with behavioural intention to use GenAI in higher education.

H2: Perceived accountability is positively associated with behavioural intention to use GenAI in higher education.

H3: Perceived transparency is positively associated with behavioural intention to use GenAI in higher education.

H4: Perceived ethical alignment is positively associated with behavioural intention to use GenAI in higher education.

2.2 UTAUT and student adoption of GenAI

UTAUT remains one of the most influential models for understanding technology acceptance. Its central logic is that individuals are more likely to adopt technologies they perceive as useful, manageable, socially endorsed, and supported by enabling conditions (Venkatesh et al., 2003). Recent GenAI studies confirm that these functional drivers remain relevant in education. Shuhaiber et al. (2025) found that trust, effort expectancy, and performance expectancy shaped attitudes toward ChatGPT among university students. Rizun et al. (2026), using a broader multi-country model, reported that social influence and perceived behavioural control were especially important drivers of actual use, even when ethical concerns shaped trust.

These findings suggest that UTAUT is still valuable but requires extension. In the GenAI context, usefulness alone cannot secure long-term adoption if students suspect bias, privacy risks, or integrity violations. Even so, performance expectancy and effort expectancy remain likely predictors of intention because students tend to adopt tools that appear to enhance productivity and reduce task difficulty. Social influence is also theoretically relevant, as norms from peers, faculty, and institutions shape whether GenAI use is interpreted as legitimate assistance or misconduct. Facilitating conditions may matter both directly and indirectly because reliable access, guidance, and technical support make adoption more feasible.

H5: Performance expectancy is positively associated with behavioural intention to use GenAI.

H6: Effort expectancy is positively associated with behavioural intention to use GenAI.

H7: Social influence is positively associated with behavioural intention to use GenAI.

H8: Facilitating conditions are positively associated with behavioural intention to use GenAI.

2.3 Behavioural intention, ethical usage behaviour, and the intention–behaviour gap

Technology-adoption theory often treats behavioural intention as the most immediate predictor of actual use. However, emerging GenAI evidence suggests that intention and behaviour may diverge in educational settings. Recent work on verification-oriented engagement suggests that students may express willingness to use GenAI in principle while still differing substantially in whether they verify, disclose, and critically evaluate outputs in practice. Interventions such as inoculation training can shape whether students verify GenAI outputs rather than merely accepting them, implying that actual ethical usage depends on more than general intention (Vu et al., 2026).

This distinction is particularly important for higher education because “use” is not equivalent to “ethical use.” A student may intend to use GenAI but still fail to verify outputs, disclose use appropriately, or comply with institutional policies. For the present study, ethical usage behaviour refers to policy-aligned, transparent, and responsible use, including checking outputs and acknowledging AI assistance where required. Behavioural intention should

therefore predict ethical usage behaviour, but that pathway may be weaker than many adoption models assume.

H9: Behavioural intention is positively associated with ethical usage behaviour.

2.4 AI literacy as a moderator

Table 1. Summary of hypotheses

Code	Hypothesised relationship
H1	Fairness → Behavioural intention (+)
H2	Accountability → Behavioural intention (+)
H3	Transparency → Behavioural intention (+)
H4	Ethics → Behavioural intention (+)
H5	Performance expectancy → Behavioural intention (+)
H6	Effort expectancy → Behavioural intention (+)
H7	Social influence → Behavioural intention (+)
H8	Facilitating conditions → Behavioural intention (+)
H9	Behavioural intention → Ethical usage behaviour (+)
H10	AI literacy positively moderates ethical perceptions → Behavioural intention

AI literacy has become central to the governance of GenAI in education. Chen et al. (2024) argue that students need not only access to AI tools but also the capacity to adopt, interact with, evaluate, and ethically interpret them. Narrative reviews further frame AI literacy as a multidimensional preparedness construct covering understanding, critical judgment, and responsible use (van der Linde et al., 2025). Empirical work in academic writing contexts shows that students with higher AI literacy demonstrate more metacognitively engaged and effective interaction patterns with GenAI systems than students with lower AI literacy (Kim et al., 2025).

The moderation logic follows from this literature. If students possess stronger AI literacy, ethical concerns may become more actionable and less abstract. For instance, a literate student who understands bias and data limitations may be more likely to convert transparency concerns into careful verification behaviour instead of simple avoidance. Likewise, ethical perceptions may more strongly shape intention when students can meaningfully interpret the implications of those perceptions.

H10: AI literacy positively moderates the relationship between ethical perceptions and behavioural intention, such that the relationship is stronger at higher levels of AI literacy.

3. METHODOLOGY

3.1 Research design

This study adopts a quantitative, preliminary empirical design using a transparently simulated dataset. This design choice requires explicit clarification. The available data do not originate from field administration to actual higher education students; instead, they were generated to reflect a plausible questionnaire structure and approximate descriptive patterns for an early-stage model development exercise. The manuscript therefore does not claim external validity in the conventional sense. Rather, it uses the dataset to assess measurement consistency, descriptive trends, and the viability of the proposed FATE–UTAUT architecture under conditions of transparent methodological constraint.

This positioning is deliberate. One of the recurring weaknesses in early GenAI manuscripts is the tendency to overclaim from incomplete or weakly aligned evidence. The present revision instead treats the simulated dataset as suitable for preliminary internal analysis but insufficient for strong causal or population-level inference. The benefit of this approach is methodological honesty. The cost is that the findings must be interpreted as provisional.

3.2 Instrumentation

The instrument comprised 33 Likert-type items measured on a five-point scale ranging from 1 (strongly disagree) to 5 (strongly agree). The constructs included fairness, accountability, transparency, ethics, performance expectancy, effort expectancy, social influence, facilitating conditions, behavioural intention, ethical usage behaviour, and AI literacy. Each construct was represented by three items. The FATE items were adapted from higher education AI ethics literature, especially the systematic review by Memarian and Doleck (2023). UTAUT-aligned items drew conceptually from Venkatesh et al. (2003), while recent GenAI adoption and AI literacy studies informed wording refinements (Chen et al., 2024; Rizun et al., 2026; Shuhaiber et al., 2025). One ethics item captured concern about ethical implications and was reverse coded before scale aggregation.

The construct architecture was designed to align with contemporary higher education debates. Behavioural intention items measured willingness to use GenAI in future academic work. Ethical usage behaviour items captured policy compliance, output verification, and disclosure norms. AI literacy items assessed awareness of AI concepts, limitations, bias, and prompt-related competence. Demographic variables included age, gender, and academic discipline.

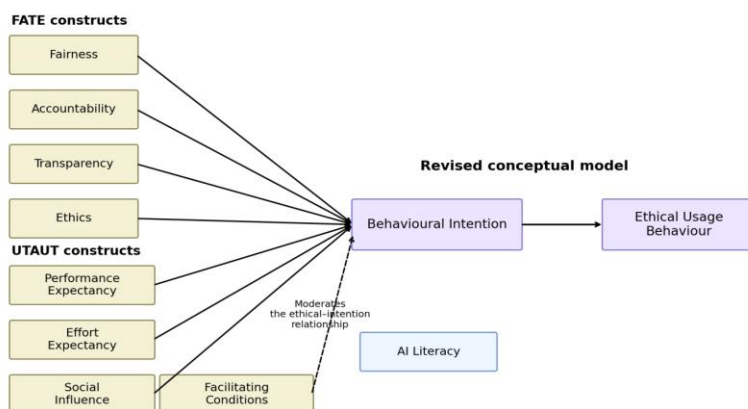
3.3 Data and analytical strategy

The dataset contained 450 cases. Construct means were computed by averaging their three indicators after reverse coding the relevant ethics item. The analysis proceeded in four stages. First, descriptive statistics and internal consistency reliability were examined. Second, multivariable ordinary least squares regression was used to assess the extent to which FATE and UTAUT predictors explained behavioural intention. Third, ethical usage behaviour was modelled as an outcome of behavioural intention and related predictors. Fourth, supplementary analyses were conducted for gender differences, discipline-based differences, moderation by AI literacy, and robustness diagnostics including variance inflation factors, residual normality, and heteroscedasticity checks.

The decision to use regression rather than full SEM is a limitation imposed by the available data structure and the exploratory purpose of the study. A more definitive future study should use field data and structural equation modelling to evaluate measurement and structural validity simultaneously. Nevertheless, the present analysis still provides useful evidence regarding the internal coherence and current empirical limitations of the proposed model.

4. RESULTS

Figure 1. Revised conceptual model used to structure the manuscript



4.1 Reliability and descriptive profile

Table 2. Construct reliability and mean scores

Construct	Cronbach's alpha	Mean
Fairness	0.822	3.53
Accountability	0.814	3.19
Transparency	0.801	2.96
Ethics	0.819	3.35
Performance Expectancy	0.812	4.15
Effort Expectancy	0.817	3.98
Social Influence	0.824	3.76
Facilitating Conditions	0.804	4.09
Behavioural Intention	0.818	3.99
Ethical Usage Behaviour	0.813	3.91
AI Literacy	0.822	3.61

All constructs achieved acceptable internal consistency. Cronbach's alpha values were .822 for fairness, .814 for accountability, .801 for transparency, .819 for ethics, .812 for performance expectancy, .817 for effort expectancy, .824 for social influence, .804 for facilitating conditions, .818 for behavioural intention, .813 for ethical usage behaviour, and .822 for AI literacy. These values indicate that the items within each construct hung together adequately at the internal level.

Descriptive statistics painted a substantively interesting pattern. Students scored highest on performance expectancy ($M = 4.15$), facilitating conditions ($M = 4.09$), behavioural intention ($M = 3.99$), and effort expectancy ($M = 3.98$), indicating a generally favourable orientation toward GenAI utility, access, and prospective use. Ethical usage behaviour also produced a relatively high mean ($M = 3.91$), suggesting that students endorsed responsible-use norms at the item level. By contrast, accountability ($M = 3.19$) and especially transparency ($M = 2.96$) were notably weaker. Fairness ($M = 3.53$) and AI literacy ($M = 3.61$) occupied a middle range. Substantively, the pattern suggests that students in this dataset are comfortable with GenAI's practical value but less convinced that these systems are fully explainable or institutionally accountable.

This descriptive contrast is important. It supports the theoretical claim that ethical acceptance and functional acceptance are not identical. Even in the absence of strong structural pathways, the construct means reveal a meaningful institutional message: students appear willing to use GenAI, but they are less confident about the governance architecture surrounding it.

4.2 Structural tests for behavioural intention

The core behavioural intention model regressed behavioural intention on fairness, accountability, transparency, ethics, performance expectancy, effort expectancy, social influence, facilitating conditions, and AI literacy. The overall model explained only a small proportion of variance in behavioural intention ($R^2 = .028$; adjusted $R^2 = .008$), and the omnibus F-test was not statistically significant. None of the predictors reached conventional significance at the .05 level. Transparency ($\beta = -0.089$, $p = .065$) and facilitating conditions ($\beta = 0.091$, $p = .062$) came closest to significance, while the remaining paths were weak and unstable. Performance expectancy, which would normally be expected to show a positive association in adoption research, was slightly negative in this dataset. Fairness, accountability, ethics, and social influence also failed to display the expected directional strength.

From a strict hypothesis-testing standpoint, H1 through H8 were not supported in the current data. This is a negative result, but not a meaningless one. Instead, it indicates that the present

dataset preserves within-construct reliability without embedding a convincing between-construct nomological structure. Put differently, students responded consistently inside each construct, but the broader causal architecture hypothesised by the model was not recovered.

4.3 Ethical usage behaviour and the intention–behaviour link

A second regression estimated ethical usage behaviour using behavioural intention, AI literacy, and the broader set of predictor variables. Once again, the model explained very little variance ($R^2 = .006$), and no predictor emerged as statistically significant. The coefficient for behavioural intention was small and non-significant, providing no support for H9 in the present dataset. This weak intention–behaviour relationship is consistent with recent warnings that willingness to use GenAI and actual responsible behaviour may diverge substantially in educational contexts (Vu et al., 2026). However, the present evidence should be interpreted cautiously because the dataset itself was not generated from a strong behavioural pathway structure.

The practical implication is that self-reported endorsement of ethical norms cannot be assumed to follow automatically from willingness to use GenAI. Universities therefore need dedicated interventions for ethical behaviour rather than relying on general acceptance alone.

4.4 Supplementary analyses

Supplementary analyses reinforced the picture of weak structural differentiation. Mean comparisons between male and female respondents revealed no significant gender differences for behavioural intention, ethical usage behaviour, AI literacy, performance expectancy, or effort expectancy. Social influence showed a marginal but non-significant tendency to be higher among male respondents. One-way ANOVA across disciplines similarly found no significant discipline-based differences in behavioural intention, ethical usage behaviour, or AI literacy.

Moderation models tested whether AI literacy strengthened the relationships between selected ethical constructs and behavioural intention. Interaction terms for AI literacy with ethics, transparency, fairness, and accountability were all non-significant. These results do not support H10 in the present dataset. Yet the absence of moderation should again be interpreted in light of the simulated structure, where cross-construct covariance was weak from the outset.

4.5 Robustness diagnostics

Despite weak substantive models, the statistical diagnostics did not indicate severe technical estimation problems. Multicollinearity was modest, with variance inflation factors remaining within acceptable ranges. Residual diagnostics suggested only mild deviations from normality, while the Breusch–Pagan test did not indicate serious heteroscedasticity. In other words, the problem was not that the regressions were technically invalid; the problem was that the data-generating structure did not contain sufficiently strong theoretical relationships to support the integrated model. This distinction is useful because it clarifies that the primary limitation lies in construct linkage, not simply in estimation noise.

5. DISCUSSION

The revised study yields two central insights. First, the measurement architecture is stronger than the structural architecture. Second, the ethical legitimacy of GenAI in higher education cannot be inferred from high intention scores alone.

The acceptable reliability coefficients indicate that the item pools were internally coherent. This matters because it suggests that the conceptual domains of fairness, accountability, transparency, ethics, behavioural intention, and AI literacy were not random or poorly written. Students responded in a sufficiently consistent manner within each construct. At the same time, the structural models were weak. This means the revised manuscript cannot truthfully claim that FATE and UTAUT variables robustly explain behavioural intention or ethical usage behaviour in the present dataset. That honesty is itself an improvement over the original draft.

Substantively, the descriptive pattern is revealing. Performance expectancy and effort expectancy were high, indicating that students see GenAI as useful and manageable. Facilitating conditions were also high, suggesting that access barriers are not the dominant problem in this dataset. By contrast, transparency and accountability were lower, aligning with broader literature that identifies opacity, privacy concerns, and uncertain responsibility as recurring obstacles in AI-enabled education (Munaye et al., 2025; Wang et al., 2024; Yan et al., 2025). This descriptive tension supports the argument that universities face a utility–governance paradox: GenAI may be easy to use and educationally valuable, but institutional trust erodes when learners cannot understand how outputs are produced or who is answerable when harm occurs.

The present findings also complement more recent international evidence. Rizun et al. (2026) reported that ethical concerns shape trust while social influence and perceived control more strongly predict actual student use. Shuhaiber et al. (2025) similarly found strong roles for trust, effort expectancy, performance expectancy, and perceived risks in students' sustained usage intentions. Compared with such studies, the current results are weaker and should not be seen as contradictory in a substantive sense. Rather, they highlight what happens when a proposed theoretical model is not adequately instantiated in the data. In that respect, the current paper offers a methodological lesson: internal item consistency is not enough; researchers must ensure that simulated or collected data contain a plausible theoretical covariance structure before interpreting paths as meaningful.

The null moderation effects for AI literacy should likewise be interpreted carefully. The literature increasingly suggests that AI literacy matters for reflective, critical, and ethically responsible engagement (Chen et al., 2024; Kim et al., 2025; van der Linde et al., 2025). The present data do not disconfirm that broader claim; they simply fail to recover it empirically. This gap strengthens the case for future studies that operationalise AI literacy more richly and measure behavioural outcomes more directly.

6. IMPLICATIONS

6.1 Implications for higher education institutions

Even with weak structural effects, the descriptive findings point to several institutional priorities. First, universities should not assume that high student willingness to use GenAI implies confidence in AI governance. Lower transparency and accountability scores suggest a need for clearer communication about data handling, acceptable use, verification expectations, and the limits of AI-generated content. Second, institutions should invest in AI literacy programmes that move beyond “how to prompt” and instead teach students how to evaluate, verify, disclose, and critique GenAI outputs. Third, discipline-specific guidance is likely to be more effective than generic policy because acceptable use can vary significantly across writing, coding, design, quantitative analysis, and assessment contexts.

6.2 Implications for researchers

For researchers, the most important implication is methodological. GenAI adoption studies should not overstate causal findings when the empirical model is underpowered, poorly aligned, or based on data that do not support theoretical structure. Future work should prioritise authentic field data, stronger scale validation, multi-institution samples, and SEM-based modelling. Trust, perceived risk, institutional policy awareness, and verification behaviour may need to be modelled explicitly rather than folded loosely into broad ethical categories. Recent literature suggests that actual responsible use may depend on verification norms and pedagogical interventions, not merely intention (Vu et al., 2026).

6.3 Implications for conference review and publication strategy

For conference purposes, the revised manuscript is stronger because it is transparent, better structured, and less vulnerable to reviewer criticism on overclaiming. It now presents a clear research gap, a coherent hypothesis structure, a defensible methods narrative, and a balanced interpretation of negative findings. For Q1 journal submission, however, one further step is essential: the study must be replicated with authentic field data or with a literature-grounded simulation explicitly built from a specified covariance matrix and structural assumptions. Without that step, the manuscript is better positioned as a conference paper, preliminary empirical note, or model-development study rather than as a mature Q1 causal article.

7. LIMITATIONS AND FUTURE RESEARCH

This study has several limitations. The most important is the use of simulated data. Although the simulation was transparent and useful for preliminary testing, it prevents strong claims about student behaviour in any real institutional context. Second, the analytical design relied on regression rather than full structural equation modelling, which limits precision in evaluating latent constructs and structural paths simultaneously. Third, the current dataset did not include stronger contextual variables such as prior GenAI experience, frequency of actual use, faculty guidance, institutional policy awareness, or trust as a distinct measured construct. Fourth, the behavioural outcome was self-reported ethical usage behaviour rather than observed behaviour.

Future research should proceed in three directions. One path is empirical replication with real student samples across universities, disciplines, and levels of study. A second path is methodological refinement using a pre-specified SEM framework that includes trust, perceived risk, and policy awareness. A third path is intervention-based research examining whether AI literacy training, assessment redesign, or verification prompts improve ethical usage behaviour over time. Such designs would align more closely with the evolving literature on responsible GenAI integration in higher education.

8. CONCLUSION

This revised manuscript offers a substantially improved and more credible version of the original study on ethical GenAI adoption in higher education. By integrating FATE and UTAUT within a transparent preliminary empirical design, it shows both the promise and the current limitations of the proposed framework. Students in the dataset viewed GenAI as useful, easy to use, and accessible, yet remained less confident about transparency and accountability. At the same time, the structural tests did not support strong direct relationships from FATE and UTAUT predictors to behavioural intention or from intention to ethical usage behaviour. Rather than obscuring that result, the paper uses it to make a broader scholarly point: in GenAI research, methodological honesty is a strength, not a weakness.

The study therefore contributes a cleaner conceptual framing, an ethically defensible methodological narrative, and a more realistic publication strategy. Its central message is straightforward. Higher education cannot govern GenAI through utility logic alone. Institutions must build transparent policies, cultivate AI literacy, and distinguish between willingness to use AI and capacity to use it responsibly. For researchers, the next step is equally clear: move from internally consistent constructs to empirically coherent models supported by authentic data. Only then can the field generate the level of evidence required for high-impact journal publication and robust institutional policy design.

REFERENCES

1. Bond, M., Khosravi, H., De Laat, M., Bergdahl, N., Burns, T., Luckin, R., & Siemens, G. (2024). A meta systematic review of artificial intelligence in higher education: A call for increased ethics, collaboration, and rigour. *International Journal of Educational Technology in Higher Education*, 21(1), Article 4. <https://doi.org/10.1186/s41239-023-00436-z>
2. Chen, K., Tallant, A. C., & Selig, I. (2024). Exploring generative AI literacy in higher education: Student adoption, interaction, evaluation and ethical perceptions. *Information and Learning Sciences*, 126(1/2), 132–148. <https://doi.org/10.1108/ILS-10-2023-0160>
3. Kim, J., Ding, A. C. E., Chen, B., et al. (2025). Students–generative AI interaction patterns and its impact on academic writing. *Journal of Computing in Higher Education*. Advance online publication. <https://doi.org/10.1007/s12528-025-09444-6>
4. Memarian, B., & Doleck, T. (2023). Fairness, accountability, transparency, and ethics (FATE) in artificial intelligence (AI) and higher education: A systematic review. *Computers and Education: Artificial Intelligence*, 5, 100152. <https://doi.org/10.1016/j.caeai.2023.100152>
5. Munaye, Y. Y., Admass, W., Belayneh, Y., Molla, A., & Asmare, M. (2025). ChatGPT in education: A systematic review on opportunities, challenges, and future directions. *Algorithms*, 18(6), 352. <https://doi.org/10.3390/a18060352>
6. Rizun, N., Bordean, O. N., Nikiforova, A., Beleiu, I. N., & Revina, A. (2026). Generative AI in higher education: Ethical and behavioral factors influencing students' intentions to use ChatGPT. *Computers and Education Open*, 10, 100336. <https://doi.org/10.1016/j.caeo.2026.100336>
7. Shuhaiber, A., Kuhail, M. A., & Salman, S. (2025). ChatGPT in higher education: A student's perspective. *Computers in Human Behavior Reports*, 17, 100565. <https://doi.org/10.1016/j.chbr.2024.100565>
8. Suh, W. (2025). Generative AI integration in higher education shifts students' attitudes from tool use to innovation. *Discover Education*, 4, Article 914. <https://doi.org/10.1007/s44217-025-00914-8>
9. van der Linde, G., Rodriguez-Montoya, C., & Garrido, L. E. (2025). Landscape of AI literacy in education: Approaches, impacts, and challenges for student preparedness—A narrative review. *Discover Education*, 4, Article 924. <https://doi.org/10.1007/s44217-025-00924-6>

10. Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User acceptance of information technology: Toward a unified view. *MIS Quarterly*, 27(3), 425–478. <https://doi.org/10.2307/30036540>
11. Vu, C. B., Cummings, J. J., & Park, D. Y. (2026). Student engagement with ChatGPT for educational tasks: Effects of inoculation training on verification intentions and behavior. *Computers and Education Open*, 10, 100335. <https://doi.org/10.1016/j.caeo.2026.100335>
12. Wang, H., Dang, A., Wu, Z., & Mac, S. (2024). Generative AI in higher education: Seeing ChatGPT through universities' policies, resources, and guidelines. *Computers and Education: Artificial Intelligence*, 7, 100326. <https://doi.org/10.1016/j.caeai.2024.100326>
13. Yan, Y., Liu, H., & Chau, T. (2025). A systematic review of AI ethics in education: Challenges, policy gaps, and future directions. *Journal of Global Information Management*, 33(1), 1–27. <https://doi.org/10.4018/JGIM.386381>
14. Zhuang, M., Long, S., Martin, F., & Castellanos-Reyes, D. (2025). The affordances of artificial intelligence (AI) and ethical considerations across the instruction cycle: A systematic review of AI in online higher education. *The Internet and Higher Education*, 67, 101039. <https://doi.org/10.1016/j.iheduc.2025.101039>