

Information Literacy and Acceptance of Generative AI among University Students and Research Scholars in India: An Extended TAM Study

Vikrant Dubey, Shweta Verma, Purnima Kumari, Virendra Kumar and Shilpi Verma

Department of Library and Information Science, Babasaheb Bhimrao Ambedkar University, Lucknow, Uttar Pradesh, India

Corresponding Author: Vikrant Dubey (vikrantdubeycup@gmail.com)

ABSTRACT

Background and Objective: The proliferation of Generative Artificial Intelligence (GenAI) tools including ChatGPT, Gemini, and Microsoft Copilot has fundamentally altered academic research workflows in Indian universities, raising critical questions about user acceptance, information literacy, and academic integrity. This study investigates how four predictors information literacy (IL), modelled as a single composite construct measured through four sub-dimensions (Source Evaluation Literacy, Search Proficiency, Ethical Use of Information, and AI Output Evaluation), together with Perceived Usefulness (PU), Perceived Credibility Risk (CR), and Academic Integrity Concern (AIC) shape scholars' attitudes toward accepting GenAI tools, within an extended Technology Acceptance Model (TAM) framework. **Materials and Methods:** A cross-sectional survey was administered to N = 137 valid respondents (after rigorous data cleaning from raw N = 140) comprising undergraduate students, postgraduate students, and PhD research scholars from central (73.2%) and state universities (26.1%) across India, spanning Humanities/Social Sciences (65.9%), Sciences/Engineering (19.6%), and Library and Information Science (13.8%) disciplines. Multiple OLS regression analyses were conducted following systematic data-quality auditing. **Results:** The model explained 66.3% of variance in GenAI acceptance attitude ($R^2 = 0.663$, Adjusted $R^2 = 0.652$, $F(4,132) = 64.79$, $p < 0.001$). All four predictors were statistically significant: Information Literacy ($\beta = 0.208$, $p = 0.005$), Perceived Usefulness ($\beta = 0.652$, $p < 0.001$), Perceived Credibility Risk ($\beta = -0.203$, $p = 0.016$), and Academic Integrity Concern ($\beta = 0.167$, $p = 0.026$). Gender, academic level, and university type did not significantly differentiate GenAI acceptance attitudes. **Conclusions:** Strengthening IL, particularly search proficiency and AI output evaluation competencies, combined with transparent communication of GenAI research utility, meaningfully promotes responsible GenAI adoption. Practical recommendations are offered for academic libraries, institutional policymakers, and GenAI tool developers.

KEYWORDS

Generative AI, technology acceptance model, information literacy, perceived usefulness, credibility risk, academic integrity, Indian universities, research scholars, library and information science

Copyright © 2026 Dubey et al. This is an open-access article distributed under the Creative Commons Attribution License, which permits unrestricted use, distribution and reproduction in any medium, provided the original work is properly cited.



INTRODUCTION

Artificial Intelligence (AI) has rapidly transformed how academic research is conducted, communicated, and evaluated. Among the most disruptive developments is the emergence of Generative AI (GenAI) tools, large language model-based platforms such as ChatGPT (OpenAI), Gemini (Google DeepMind), and Microsoft Copilot, which have moved from experimental tools to everyday instruments in university research workflows within a remarkably short period^{1,2}. In India, this adoption has been particularly rapid: A country with one of the world's largest higher education systems and a growing cadre of research scholars is grappling simultaneously with the opportunity and the risk that GenAI presents for academic work.

Despite this urgency, the cognitive and literacy-based factors that determine whether scholars embrace or resist GenAI tools remain underexplored in the Indian context. The Technology Acceptance Model (TAM; Davis)³ has long been the dominant framework for studying user adoption of information systems. TAM centres on Perceived Usefulness (PU) and Perceived Ease of Use (PEOU) as the principal determinants of adoption intention and behaviour. However, scholars have argued that TAM in its original form is insufficient for capturing the nuanced cognitive, ethical, and risk-related dimensions of AI-mediated academic work^{4,5}. There is a growing consensus that information literacy, the competency to locate, evaluate, and ethically use information, plays a crucial mediating role in how scholars perceive and engage with AI-generated content⁶.

This study advances three arguments. First, information literacy is not merely a desirable graduate attribute but an active cognitive resource that shapes technology acceptance attitudes in the GenAI era. Second, credibility risk concerns that GenAI produces hallucinated, fabricated, or unreliable scholarly content constitute a distinct and measurable barrier to adoption that TAM's original constructs do not capture. Third, academic integrity concern operates as a double-edged variable: Rather than deterring adoption, integrity-aware scholars may engage with GenAI tools more deliberately, strategically, and responsibly.

Taken together, these arguments expose a clear gap in the existing literature. Prior TAM-based studies of AI and chatbot adoption in education have extended the model with constructs such as trust, perceived AI capability, performance expectancy, and general technological literacy, but they have treated literacy either as a single undifferentiated trait or as a peripheral control rather than as a structured, multi-dimensional information-literacy construct. Parallel work in information science has richly conceptualised information and algorithmic literacy, yet has rarely embedded these competencies within a formal technology-acceptance model or tested them as direct predictors of GenAI acceptance. Furthermore, credibility risk and academic integrity concern, although widely discussed in commentary on GenAI, have seldom been operationalised and modelled simultaneously alongside information literacy and perceived usefulness within a single empirical framework, and rarely in the Indian higher-education context. The present study addresses this gap by integrating a four-dimensional conception of information literacy (source evaluation, search proficiency, ethical use, and AI output evaluation) with perceived usefulness, perceived credibility risk, and academic integrity concern into one extended TAM model, and by testing it on rigorously audited data from Indian university students and research scholars. In doing so, it offers one of the first model-based accounts of how cognitive competencies and risk perceptions jointly shape GenAI acceptance in this setting.

Study addresses three research questions:

- **RQ1:** How do the four dimensions of information literacy, source evaluation, search proficiency, ethical use of information, and AI output evaluation individually and collectively influence university students' and research scholars' attitudes toward accepting GenAI tools?

- **RQ2:** What is the relative explanatory contribution of perceived usefulness, perceived credibility risk, and academic integrity concern in predicting GenAI acceptance attitudes?
- **RQ3:** Do demographic characteristics, gender, academic level, discipline, and university type significantly differentiate GenAI acceptance attitudes?

The study is grounded in clean, rigorously audited survey data from N = 137 respondents spanning undergraduate, postgraduate, and doctoral levels at central and state universities across India. The extended TAM model explains 66.3% of variance in GenAI acceptance attitude a result that compares favourably with prior TAM studies of standard two-factor TAM applications in similar domains. The findings carry direct implications for academic library instruction, institutional GenAI policy, and tool design.

THEORETICAL FRAMEWORK AND HYPOTHESES

Technology acceptance model: Extensions and limitations: Davis³ developed TAM to explain and predict users' acceptance of information technology. The model posits that Perceived Usefulness (PU), the degree to which a technology is believed to enhance job performance, and Perceived Ease of Use (PEOU) directly determine attitude toward using technology, which in turn drives behavioural intention and actual use. TAM has accumulated an impressive body of empirical support across diverse technological contexts⁷.

The TAM's parsimony, often cited as a strength, becomes a limitation in domains where domain-specific cognitive competencies and risk perceptions mediate adoption. Several scholars have extended TAM to incorporate constructs such as trust⁸, perceived AI capability⁹, performance expectancy⁵, and technological literacy¹⁰. In the GenAI domain specifically, Rabiul Awal and Enamul Haque¹¹ demonstrated that social cognitive factors enhance TAM's predictive validity. Yet none of these extensions systematically integrates information literacy as a structured multi-dimensional predictor a gap this study addresses.

Information literacy in the age of generative AI: The American Library Association¹² defines information literacy as a set of integrated abilities encompassing the reflective discovery of information, understanding how information is produced and valued, and the ethical use of information. In the GenAI era, these competencies take on heightened significance. AI-generated text can appear credible, well-structured, and authoritatively cited while containing fabricated references, outdated facts, or subtle factual errors, phenomena termed 'hallucinations'¹³.

This study operationalises information literacy through four domain-specific dimensions. Source Evaluation Literacy (SE) concerns the ability to critically assess the provenance, authority, and reliability of AI-generated outputs. Search Proficiency (SP) encompasses the competency to construct effective prompts and queries across traditional databases and AI platforms. Ethical Use of Information (EU) covers understanding of attribution norms, academic integrity boundaries, and responsible parameters for AI assistance. AI Output Evaluation (AE) refers specifically to the ability to detect hallucinations, cross-verify AI-generated claims, and assess the scholarly quality of AI outputs. Archambault *et al.*⁶ have argued that AE and prompt literacy are now indispensable dimensions of IL instruction in academic library settings.

Perceived credibility risk: Perceived risk theory¹⁴ describes users' assessment of the potential negative outcomes of adopting a technology or service. In the context of scholarly research, credibility risk takes a particularly acute form: The possibility that AI-generated content undermines research integrity through false citations, fabricated statistics, or misleading synthesis. Kaushal and Yadav¹⁵ found that users of academic AI tools were acutely concerned about the reliability of AI responses for research-critical tasks. Wang *et al.*¹⁶ demonstrated that perceived risk significantly attenuates adoption intentions for healthcare chatbots, a finding that the present study extends to the academic GenAI context.

Academic integrity concern: Academic integrity concern (AIC) is a construct specific to educational contexts: the degree to which users worry about policy violations, institutional detection, and reputational consequences of AI-assisted work submission¹⁷. The relationship between AIC and GenAI acceptance is theoretically ambiguous. On one hand, high integrity concerns could deter adoption if scholars fear institutional sanctions. On the other hand, scholars who are well-informed about integrity norms may engage with GenAI tools more strategically and within permissible boundaries, treating integrity awareness as a framework for responsible rather than avoidant use. This study tests this ambiguity empirically. Because these two mechanisms predict opposite signs, the theoretical literature does not support a single directional expectation for the effect of academic integrity concern. Accordingly, H4 was deliberately formulated a priori as a non-directional (two-tailed) hypothesis, predicting that academic integrity concern has a significant effect on GenAI acceptance attitude without specifying its sign in advance. This non-directional framing reflects the genuinely competing predictions described above (deterrence versus responsible-engagement) and allows the data to adjudicate between them, rather than imposing a directional assumption that the theory cannot justify.

Demographic moderators: Several studies have examined whether demographic characteristics moderate technology acceptance. Dogruel *et al.*¹⁸ found that gender and education level influenced algorithm literacy. Trepte *et al.*¹⁹ demonstrated that demographic attributes shape digital literacy profiles. This study extends this line of inquiry to the GenAI context by examining whether gender, academic level, discipline, and university type significantly differentiate GenAI acceptance attitudes among Indian university scholars.

Hypotheses: Based on the above literature, the following hypotheses are proposed as follows:

- **H1:** Information literacy (composite of SE, SP, EU, AE) has a significant positive effect on attitude toward accepting GenAI tools for academic research
- **H2:** Perceived usefulness has a significant positive effect on attitude toward accepting GenAI tools
- **H3:** Perceived credibility risk has a significant negative effect on attitude toward accepting GenAI tools
- **H4:** Academic integrity concern has a significant effect on attitude toward accepting GenAI tools
- **H5:** Demographic characteristics (gender, academic level, discipline, university type) significantly differentiate GenAI acceptance attitudes

RESEARCH METHODOLOGY

Study area and duration: The study was conducted among undergraduate students, postgraduate students, and PhD research scholars enrolled in central and state universities across India. Data collection was carried out from 11 December 2025 to 10 January 2026 using an online questionnaire.

Research design: Based on the above theoretical framework and hypotheses, the conceptual model of the study is presented in Fig. 1.

This study adopts a quantitative, cross-sectional survey design grounded in the positivist research paradigm. A structured self-administered questionnaire was used to collect primary data, and Ordinary Least Squares (OLS) multiple regression analysis was employed to test the hypothesised relationships. The study is theoretically anchored in an extended TAM framework that incorporates information literacy constructs and risk perceptions as additional predictors.

Population, sampling, and data collection: The target population comprised undergraduate students, postgraduate students, and PhD research scholars at central and state universities across India who had prior experience using at least one GenAI tool (ChatGPT, Gemini, Copilot, or equivalent) for academic

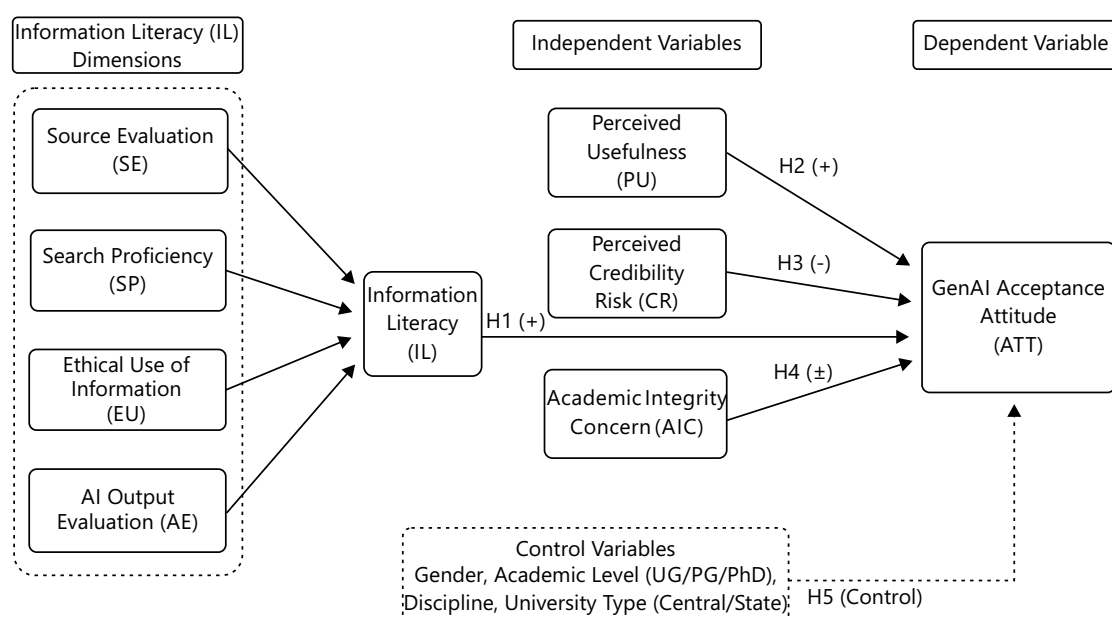


Fig. 1: Conceptual framework of GenAI acceptance: An extended TAM

Solid arrows indicate hypothesized direct effect, Dashed arrow control effect, (+) = Positive, (-) = Negative relationship and (±) = Non-directional relationship

research purposes. Purposive sampling was employed to ensure that respondents had relevant experience with the phenomenon under investigation. Data were collected anonymously through a structured online questionnaire that was circulated via academic social-networking platforms (including LinkedIn, Facebook, and WhatsApp research groups) and researcher peer networks, thereby reaching respondents from different regions of India. The data-collection period extended from 11 December 2025 to 10 January 2026.

The initial dataset comprised $N = 140$ entries. Before analysis, a systematic row-by-row data audit was conducted. Two records were excluded: Row 139 contained numeric junk values in multiple demographic fields (Gender coded as '71'; Academic Level as '45'), indicating a corrupted or test entry; Row 140 contained entirely missing Likert responses. Additionally, several Attitude toward Accepting GenAI items were found to contain backtick-prefixed entries ('N') a data-entry artefact arising from keyboard input, recoded to 'N' (Neutral = 3) before analysis. After cleaning, the final valid sample for demographic analysis comprised $N = 138$; construct-level analyses used $N = 137$ (one additional respondent had missing Likert data across all constructs). To summarise the analytical flow in one place: of the 140 raw responses collected, 2 were removed during data cleaning (Rows 139 and 140), leaving 138 valid cases. All 138 cases carried complete demographic information and were therefore retained for the descriptive sample-profile analysis in sample characteristics (Table 2). One of these 138 respondents had missing values across all Likert-based construct items and so could not be assigned construct scores; this case was excluded listwise from every analysis that relies on construct scores, namely the reliability, correlation, regression, and demographic group-difference analyses, which are accordingly based on $N = 137$ (Table 3-7). The two sample sizes therefore differ only because the demographic profile in Table 2 requires only demographic fields, whereas all construct-based analyses (including the group comparisons of attitude scores in Table 7) additionally require complete Likert responses. This distinction is reiterated in the relevant table notes for readability.

Measurement instrument: The questionnaire comprised four sections: (A) Demographic information; (B) Information literacy items; (C) TAM and risk perception items; and (D) Attitude toward GenAI acceptance items. All substantive items were measured on a five-point Likert scale anchored at 1 (Strongly Disagree) to 5 (Strongly Agree). Table 1 summarises the eight constructs, item counts, Cronbach's alpha reliability values, and primary source references.

Table 1: Measurement constructs, item count, reliability, and sources (N = 137)

Construct	Code	Items	A	Primary Source
Source Evaluation Literacy	SE	5	0.832	ALA ¹² and Archambault <i>et al.</i> ⁶
Search Proficiency	SP	5	0.860	Nikou <i>et al.</i> ²⁰
Ethical Use of Information	EU	5	0.911	Cetindamar <i>et al.</i> ²¹
AI Output Evaluation	AE	5	0.928	Shin <i>et al.</i> ²²
Perceived Usefulness	PU	5	0.905	Davis ³ and Safadel <i>et al.</i> ²³
Perceived Credibility Risk	CR	5	0.800	Featherman and Pavlou ¹⁴
Academic Integrity Concern	AIC	5	0.862	Stieglitz <i>et al.</i> ¹⁷
Attitude Toward Accepting GenAI	ATT	5	0.874	DeLone and McLean ²⁴

Data analysis strategy: Data were analysed in Python 3.11 using pandas, scipy, and numpy libraries. The analysis pipeline comprised six sequential steps: (1) Full column-level audit of all unique values to identify encoding artefacts, junk entries, and missing data; (2) Removal of invalid rows; (3) Correction of backtick-prefixed Likert values; (4) Likert encoding (SD=1, D=2, N=3, A=4, SA=5) and construct score computation as item means; (5) Internal reliability assessment using Cronbach's alpha ($\alpha \geq 0.70$ threshold²⁵); (6) Pearson bivariate correlations; (7) OLS multiple linear regression with composite IL, PU, CR, and AIC as predictors of ATT; and (8) Supplementary group difference tests independent samples t-tests for gender and university type, and one-way ANOVA for academic level to address RQ3 and H5. Information Literacy was modelled as a single composite predictor (the mean of its four sub-dimensions: SE, SP, EU, and AE) rather than as four separate regressors, for three reasons. First, this is theoretically consistent with the study's conceptualisation of information literacy as one higher-order competency expressed through complementary facets, in line with the integrated framework of the American Library Association (2016). Second, the four sub-dimensions were strongly inter-correlated ($r = 0.65-0.81$), so entering them simultaneously would introduce substantial multicollinearity, inflate standard errors, and render the individual coefficients unstable and difficult to interpret. Third, aggregating the conceptually related items into one reliable composite improves parsimony and statistical power relative to the sample size, reducing the risk of over-fitting in a model already containing PU, CR, and AIC. The bivariate contributions of the individual sub-dimensions are nonetheless reported in Section 4.4 to preserve dimensional insight.

Ethical considerations: This study involved human participants through a voluntary, anonymous questionnaire survey and posed no more than minimal risk to respondents. Participation was entirely voluntary, and no personally identifying information (such as name, contact details, or institutional identifiers) was collected. Informed consent was obtained from every participant by means of a consent statement presented at the beginning of the online questionnaire, which explained the purpose of the study, the voluntary and anonymous nature of participation, the intended academic use of the data, and the respondent's right to withdraw before submission; participants could proceed to the survey items only after indicating their consent. Given the anonymous, non-interventional, and minimal-risk nature of the survey, formal clearance from an institutional review board or ethics committee was not separately required for this study; nonetheless, the research was conducted in accordance with the ethical principles of the Declaration of Helsinki and the data-protection norms of the participating institutions.

RESULTS

Sample characteristics: The final sample for demographic analysis comprised N = 138 respondents. Gender was nearly equally distributed: 70 women (50.7%), 67 men (48.6%), and 1 missing (0.7%). The dominant age group was 26-30 years ($n = 88$, 63.8%), followed by 22-25 years ($n = 28$, 20.3%), 31-35 years ($n = 10$, 7.2%), above 35 years ($n = 7$, 5.1%), and 18-21 years ($n = 4$, 2.9%). Academic level composition showed PhD/Research Scholars as the largest group ($n = 58$, 42.0%), followed by Undergraduates ($n = 45$, 32.6%), and Postgraduate students ($n = 34$, 24.6%), with one missing response (0.7%).

Table 2: Sample profile (N = 138)

Characteristics	Category	N	%
Gender	Man	67	48.6
	Woman	70	50.7
	Missing	1	0.7
Age group	18-21	4	2.9
	22-25	28	20.3
	26-30	88	63.8
	31-35	10	7.2
	Above 35	7	5.1
	Missing	1	0.7
Academic level	Undergraduate	45	32.6
	Postgraduate	34	24.6
	PhD/Research Scholar	58	42.0
	Missing	1	0.7
Discipline	Humanities/Social Sciences	91	65.9
	Sciences/Engineering	27	19.6
	Library and Information Science	19	13.8
	Missing	1	0.7
University type	Central University	101	73.2
	State University	36	26.1
	Missing	1	0.7
Library usage	Daily	68	49.3
	Weekly	33	23.9
	Seasonal	26	18.8
	Monthly	7	5.1
	Never	3	2.2
	Missing	1	0.7
GenAI usage frequency	Daily	90	65.2
	Weekly	30	21.7
	Monthly	14	10.1
	Seasonal	3	2.2
	Missing	1	0.7

Demographic profile is based on all 138 valid respondents. Construct-based analyses (Tables 3-7) use N = 137, because one respondent had missing values on all Likert items and could not be assigned construct scores

Table 3: Reliability analysis (N = 137)

Construct	Code	No. Items	Cronbach's α	Interpretation
Source Evaluation Literacy	SE	5	0.832	Good
Search Proficiency	SP	5	0.860	Good
Ethical Use of Information	EU	5	0.911	Excellent
AI Output Evaluation	AE	5	0.928	Excellent
Perceived Usefulness	PU	5	0.905	Excellent
Perceived Credibility Risk	CR	5	0.800	Good
Academic Integrity Concern	AIC	5	0.862	Good
Attitude Toward Accepting GenAI	ATT	5	0.874	Good

Discipline representation across three categories was: Humanities and Social Sciences (n = 91, 65.9%), Sciences and Engineering (n = 27, 19.6%), and Library and Information Science (n = 19, 13.8%), with one missing response (0.7%). University type was distributed across Central University (n = 101, 73.2%) and State University (n = 36, 26.1%), with one missing (0.7%). Library usage was high: 68 respondents (49.3%) used the library daily. GenAI tool usage was very frequent: 90 respondents (65.2%) used GenAI tools daily and 30 (21.7%) weekly, indicating a highly experienced sample well-positioned to report informed acceptance attitudes. Full sample details appear in Table 2.

Reliability analysis: All eight constructs demonstrated satisfactory to excellent internal consistency, exceeding the threshold of $\alpha \geq 0.70^{25}$. The AI Output Evaluation recorded the highest alpha ($\alpha = 0.928$), followed by Ethical Use of Information ($\alpha = 0.911$) and Perceived Usefulness ($\alpha = 0.905$). The dependent variable, Attitude toward Accepting GenAI, yielded $\alpha = 0.874$. The lowest alpha was for Perceived Credibility Risk ($\alpha = 0.800$), which nonetheless exceeds acceptable thresholds. Full reliability statistics are presented in Table 3.

Table 4: Descriptive statistics (N = 137)

Construct	N	Mean	SD	Median	Min	Max
Source Evaluation Literacy (SE)	137	3.831	0.631	4.000	2.20	5.00
Search Proficiency (SP)	137	3.867	0.669	4.000	1.40	5.00
Ethical Use of Information (EU)	137	3.930	0.776	4.000	1.80	5.00
AI Output Evaluation (AE)	137	3.784	0.823	4.000	1.20	5.00
Perceived Usefulness (PU)	137	3.799	0.748	4.000	1.60	5.00
Perceived Credibility Risk (CR)	137	3.715	0.658	3.800	2.60	5.00
Academic Integrity Concern (AIC)	137	3.635	0.686	3.600	2.00	5.00
Attitude Toward Accepting GenAI (ATT)	137	3.686	0.702	3.600	2.00	5.00

Table 5: Pearson correlation matrix

	SE	SP	EU	AE	PU	CR	AIC	ATT
SE	1.000	0.651***	0.751***	0.812***	0.510***	0.283***	0.240**	0.434***
SP	-	1.000	0.795***	0.721***	0.775***	0.520***	0.358***	0.756***
EU	-	-	1.000	0.807***	0.545***	0.534***	0.380***	0.476***
AE	-	-	-	1.000	0.537***	0.366***	0.271**	0.556***
PU	-	-	-	-	1.000	0.493***	0.407***	0.793***
CR	-	-	-	-	-	1.000	0.712***	0.360***
AIC	-	-	-	-	-	-	1.000	0.378***
ATT	-	-	-	-	-	-	-	

1.000***p<0.001, **p<0.01, N = 137 and Lower triangle omitted for readability

Table 6: Multiple regression results (DV: Attitude toward accepting GenAI)

Predictor	Unstd. β	Std. Error	t-Statistic	p-value	95% CI	Decision
Intercept	0.556	0.250	2.226	0.028*	[0.062, 1.051]	-
Information Literacy (IL)	0.208	0.073	2.842	0.005**	[0.063, 0.352]	H1 Supported
Perceived Usefulness (PU)	0.652	0.065	9.985	<0.001***	[0.523, 0.781]	H2 Supported
Perceived Credibility Risk (CR)	-0.203	0.083	-2.449	0.016*	[-0.367, -0.039]	H3 Supported
Academic Integrity Concern (AIC)	0.167	0.074	2.256	0.026*	[0.021, 0.313]	H4 Supported

*p<0.05, **p<0.01, ***p<0.001. $R^2 = 0.663$, Adjusted $R^2 = 0.652$, $F(4,132) = 64.79$, $p<0.001$. N = 137. 95% CI = 95% Confidence Interval and H4: AIC was statistically significant formulated as a non-directional hypothesis; the observed effect was positive

Descriptive statistics: Table 4 presents descriptive statistics for all eight constructs (N = 137). Ethical Use of Information recorded the highest mean (M = 3.930, SD = 0.776), indicating strong normative awareness of ethical information use among respondents. Search Proficiency (M = 3.867, SD = 0.669) and Perceived Usefulness (M = 3.799, SD = 0.748) were also above the scale midpoint of 3.0. Attitude toward Accepting GenAI yielded a mean of 3.686 (SD = 0.702), indicating a moderately positive overall acceptance disposition. Academic Integrity Concern recorded the lowest mean (M = 3.635, SD = 0.686). Notably, all eight construct means exceeded the scale midpoint, suggesting generally positive orientations across all measured dimensions.

Bivariate correlation analysis: Table 5 presents Pearson bivariate correlations among all constructs. Perceived Usefulness showed the strongest bivariate correlation with ATT ($r = 0.793$, $p<0.001$), followed closely by Search Proficiency ($r = 0.756$, $p<0.001$). These two predictors achieved nearly identical bivariate strength, underscoring that practical research-utility perceptions and practical IL competencies are co-equally important antecedents of acceptance. AI Output Evaluation ($r = 0.556$, $p<0.001$), Ethical Use of Information ($r = 0.476$, $p<0.001$), and Source Evaluation Literacy ($r = 0.434$, $p<0.001$) also demonstrated significant positive associations with ATT. The positive bivariate correlations of CR ($r = 0.360$, $p<0.001$) and AIC ($r = 0.378$, $p<0.001$) with ATT appear paradoxical, but risk-aware scholars still lean toward acceptance. This pattern is resolved in the multivariate model where, controlling for IL and PU, CR's negative suppressive role becomes apparent, illustrating the indispensability of multivariate over bivariate analysis.

Multiple regression analysis: Testing H1-H4: Full regression results appear in Table 6. The regression analysis showed that all four predictors significantly influenced the dependent variable. Information Literacy had a significant positive effect ($\beta = 0.208$, $p = 0.005$), supporting H1. Perceived Usefulness was

Table 7: Demographic group difference tests (N = 137)

Comparison	Test	Statistic	p-value	Decision
Gender (Man vs. Woman)	Independent t-test	$t = 0.590$	0.556	H5 Not Supported
Academic Level (UG/PG/PhD)	One-way ANOVA	$F(2,134) = 2.062$	0.131	H5 Not Supported
University Type (Central vs. State)	Independent t-test	$t = 1.079$	0.282	H5 Not Supported

ATT: Attitude toward Accepting GenAI Tools. None of the demographic comparisons reached statistical significance ($p > .05$). These group comparisons are based on respondents' attitude (ATT) construct scores and therefore use $N = 137$ (the one respondent with missing Likert data across all constructs is excluded), consistent with the construct-level analyses in Tables 3-6

Table 8: Summary of all hypothesis tests

Hypothesis	Relationship	Result	Key Statistic
H1	IL $\rightarrow$ ATT (positive)	Supported	$\beta = 0.208, p = 0.005$
H2	PU $\rightarrow$ ATT (positive)	Supported	$\beta = 0.652, p < 0.001$
H3	CR $\rightarrow$ ATT (negative)	Supported	$\beta = -0.203, p = 0.016$
H4	AIC $\rightarrow$ ATT (effect)	Supported [†]	$\beta = +0.167, p = 0.026$
H5	Demographics $\rightarrow$ ATT	Not supported	All $p > 0.05$

[†]Significant positive effect

the strongest positive predictor ($\beta = 0.652, p < 0.001$), supporting H2. Perceived Credibility Risk had a significant negative effect ($\beta = -0.203, p = 0.016$), supporting H3. Academic Integrity Concern also positively predicted the outcome ($\beta = 0.167, p = 0.026$), supporting H4. Overall, all proposed hypotheses were supported. This indicates that the four predictors collectively explain 66.3% of the variance in GenAI acceptance attitude, a strong result relative to comparable TAM studies in the educational technology domain (cf.^{5,11}). All four predictors were statistically significant (Table 8).

Demographic group differences: Testing H5: To test H5, supplementary group difference analyses were conducted. Results are summarised in Table 7. The independent-samples t-test indicated that there was no significant difference in the study variable between men and women ($t = 0.590, p = 0.556$). Similarly, the one-way ANOVA revealed no significant differences among undergraduate (UG), postgraduate (PG), and PhD students ($F(2,134) = 2.062, p = 0.131$). Furthermore, the independent-samples t-test showed no significant difference between students from central and state universities ($t = 1.079, p = 0.282$). Since all p-values were greater than .05, no statistically significant demographic differences were observed across gender, academic level, or university type. Therefore, H5 was not supported.

Gender did not significantly differentiate ATT scores (Men: $M = 3.722, SD = 0.860$; Women: $M = 3.651, SD = 0.512; t = 0.590, p = 0.556$). Academic level similarly showed no significant difference across undergraduate ($M = 3.698, SD = 0.873$), postgraduate ($M = 3.488, SD = 0.579$), and doctoral ($M = 3.793, SD = 0.598$) respondents ($F(2,134) = 2.062, p = 0.131$). University type did not differentiate ATT: Central University respondents ($M = 3.725, SD = 0.678$) and State University respondents ($M = 3.578, SD = 0.765$) did not differ significantly ($t = 1.079, p = 0.282$). These results consistently fail to support H5, suggesting that GenAI acceptance attitudes in this sample are driven by cognitive and perceptual factors (IL, PU, CR) rather than demographic characteristics.

DISCUSSION

Information literacy as a cognitive driver of GenAI acceptance: The significant positive effect of composite Information Literacy on GenAI acceptance ($\beta = 0.208, p = 0.005$) is the study's central theoretical contribution. By integrating IL into the TAM framework, this study demonstrates that cognitive competency is a meaningful predictor of AI technology adoption, a dimension overlooked in earlier TAM extensions. Scholars who are proficient in source evaluation, search construction, ethical use, and AI output verification approach GenAI tools with greater self-efficacy and strategic intentionality, translating into more favourable acceptance attitudes.

Among the four IL sub-dimensions, Search Proficiency showed the strongest bivariate association with ATT ($r = 0.756$), nearly matching Perceived Usefulness ($r = 0.793$). This finding is instructive: Scholars who know how to construct effective GenAI prompts analogous to the database search skills that information

literacy has long cultivated extract greater utility from these tools, reinforcing positive acceptance perceptions. This resonates with Archambault *et al.*⁶, who argue that prompt literacy is now an indispensable extension of traditional information literacy instruction in academic libraries. AI Output Evaluation ($r = 0.556$) similarly showed a strong association, suggesting that scholars who can detect and correct AI hallucinations are more confident and positive in their engagement with these tools.

Perceived usefulness as the dominant predictor: Perceived Usefulness was the strongest predictor in the regression model ($\beta = 0.652$, $p < 0.001$), consistent with foundational TAM theory³ and recent AI adoption studies^{5,11}. This finding reflects the pragmatic orientation of this sample: 65.2% use GenAI tools daily, and 42.0% are doctoral scholars engaged in intensive research activity. Frequent, experienced users have accumulated direct evidence of GenAI utility accelerating literature review, structuring arguments, and retrieving information and this lived utility experience strongly conditions positive acceptance attitudes. The practical implication is clear: Institutional adoption strategies should foreground concrete demonstrations of research-specific usefulness rather than generic technology promotion. Situating these results within the broader international evidence base reinforces their interpretation while also revealing context-specific nuances. The dominance of perceived usefulness echoes findings from South Asian and Middle Eastern higher-education samples, where Rabiul Awal and Enamul Haque¹¹ and Pillai *et al.*⁵ similarly reported usefulness and performance expectancy as the leading determinants of AI-chatbot acceptance. At the same time, the substantial and independent contribution of information literacy observed here goes beyond most prior TAM extensions, which have tended to privilege trust or general technological literacy; this suggests that, in research-intensive contexts, domain-specific competencies may matter more than the broad digital-literacy measures used in European algorithm-literacy studies. The negative effect of credibility risk is consistent with international work on risk in AI and chatbot adoption, including healthcare-chatbot evidence that perceived risk dampens adoption intentions, indicating that credibility concerns operate similarly across cultural and disciplinary boundaries. Read against this wider literature, the present findings both corroborate the cross-cultural primacy of perceived usefulness and extend it by demonstrating that structured information literacy is a distinct, competing driver of GenAI acceptance.

Perceived credibility risk as a significant adoption barrier: The significant negative effect of Perceived Credibility Risk ($\beta = -0.203$, $p = 0.016$) confirms H3 and extends perceived risk theory¹⁴ to the academic GenAI context. Despite overall positive acceptance attitudes, concern about AI hallucinations and unreliable scholarly outputs meaningfully attenuates acceptance when other factors are statistically controlled. This finding has direct practical implications: GenAI tool developers and academic institutions must invest in transparency mechanisms, citation verification features, output confidence indicators, and hallucination warnings to reduce credibility risk perceptions among academically sophisticated users.

The apparent paradox that CR correlated positively with ATT in bivariate analysis ($r = 0.360$) but negatively in the multivariate model is theoretically important. It demonstrates that credibility-conscious scholars are not deterred from adoption in absolute terms; rather, their acceptance is moderated once the dominant utility effect (PU) and literacy effects (IL) are accounted for. This pattern underscores that bivariate correlation alone would have led to incorrect conclusions about the CR-ATT relationship.

Academic integrity concern: Responsibility rather than resistance: Academic Integrity Concern was statistically significant but in an unexpected positive direction ($\beta = 0.167$, $p = 0.026$). This finding challenges the assumption that integrity concerns deter GenAI adoption. One theoretically coherent interpretation is that integrity-aware scholars engage with GenAI tools more deliberately and within self-imposed ethical boundaries, using these tools as legitimate research aids rather than as shortcuts for unethical work. This is consistent with Noguera-Vivo and del Mar Grandío-Pérez²⁶, who found that algorithmic awareness promotes responsible rather than avoidant technology engagement. Another

explanation may be that the absence of a clear institutional GenAI policy at many Indian universities means integrity concern does not yet operationalise as a concrete deterrent. Future research should examine whether this positive effect is moderated by institutional policy clarity.

Demographic invariance of GenAI acceptance attitudes: The non-significant demographic group differences (all $p > 0.05$ for gender, academic level, and university type) are a substantively important finding. They suggest that in this sample, GenAI acceptance attitudes are not stratified by who the scholar is but by what they know and perceive about GenAI tools. This has egalitarian implications: IL enhancement programmes and effective usefulness communication should be equally accessible and beneficial across gender, academic level, discipline, and university type. The non-significance of the academic level ANOVA ($F = 2.062$, $p = 0.131$) is particularly notable: undergraduate students ($M = 3.698$) and doctoral scholars ($M = 3.793$) showed remarkably similar GenAI acceptance attitudes, suggesting that positive GenAI orientations are not confined to more advanced researchers^{25,26}.

PRACTICAL IMPLICATIONS

For academic libraries: Academic libraries occupy the pivotal instructional position for IL development. This study's findings suggest three priority areas. First, libraries should develop GenAI-specific prompt literacy workshops that build on established database search instruction competencies, given that Search Proficiency showed the strongest IL-ATT correlation ($r = 0.756$). Second, AI Output Evaluation training covering hallucination detection, AI citation verification, and output quality assessment should be integrated into research skills programmes at all academic levels. Third, since the study found no significant difference across academic levels, these programmes should be designed for undergraduates and doctoral scholars alike. To make these recommendations concrete, the prompt-literacy workshops could take the form of a short, hands-on module in which students iteratively refine a research query, compare the outputs of two GenAI tools against a library database search, and document which prompt formulations yield the most accurate and verifiable results. AI Output Evaluation training could be operationalised through a recurring "verify-the-AI" exercise embedded in dissertation and thesis support: students are given a GenAI-generated literature summary with deliberately planted hallucinated citations and asked to detect, trace, and correct them using the library's discovery tools. Libraries could further institutionalise these competencies by adding a one-credit "Responsible GenAI for Research" micro-course to the existing information-literacy curriculum, and by creating a shared online repository of vetted prompt templates and tool-evaluation checklists that subject librarians maintain for each discipline. Such measures translate the abstract competencies of source evaluation, search proficiency, ethical use, and AI output evaluation into routine, assessable activities.

For institutional policymakers: The positive effect of Academic Integrity Concern on acceptance attitudes suggests that integrity awareness, when properly channelled through clear institutional policy, promotes responsible rather than avoidant GenAI engagement. Universities should develop explicit, discipline-specific GenAI use policies that distinguish permissible research assistance from prohibited ghostwriting. Given the balanced representation of Central (73.2%) and State (26.1%) universities in this sample, policy guidance should be disseminated through both central university networks and state-level higher education councils.

For GenAI tool developers: Perceived Usefulness is the dominant adoption driver ($\beta = 0.652$). Developers should invest in demonstrating concrete research-usefulness through library database integrations, academic citation verification features, and research-workflow-specific onboarding. Reducing perceived credibility risk through transparent confidence indicators and hallucination warnings would further enhance adoption among academically sophisticated users.

LIMITATIONS AND FUTURE RESEARCH DIRECTIONS

This study makes a meaningful empirical contribution but acknowledges several limitations that should inform future research.

First, although the sample spans three discipline groups, two university types, and respondents from different regions of India, it was recruited through purposive online sampling and is not a probability sample, which limits the statistical generalisability of the findings. The Humanities/Social Sciences group also accounts for 65.9% of respondents, constraining conclusions about STEM-dominant institutional contexts. A further sampling-related limitation is the potential for self-selection bias: Because purposive sampling was combined with a voluntary online questionnaire distributed through academic social-networking platforms, respondents who chose to participate may have been systematically more interested in, more experienced with, or more favourably disposed towards GenAI tools than non-respondents. This may partly explain the generally positive construct means and could inflate observed acceptance levels, so the descriptive levels reported here should be interpreted with appropriate caution. Future studies should deliberately stratify samples across disciplines, regions, and institutional types, and employ probability-based or multi-wave recruitment designs to reduce self-selection bias.

Second, the OLS regression approach, while methodologically appropriate for the sample size and research questions, does not allow for structural modelling of indirect or mediated pathways. Future research should employ Partial Least Squares Structural Equation Modelling (PLS-SEM), to test mediation effects, for example, whether IL mediates the effect of PU on ATT and generate composite reliability (CR) and average variance extracted (AVE) statistics for convergent and discriminant validity assessment.

Third, the study measures attitude toward GenAI acceptance rather than actual adoption behaviour. The attitude-behaviour gap is well-documented¹⁷: Positive attitudes do not guarantee sustained or responsible use. Future research should incorporate behavioural intention, frequency of use, and quality of GenAI integration into research outputs as outcome variables.

Fourth, the cross-sectional design precludes causal inference. Longitudinal designs would better capture how GenAI acceptance attitudes evolve as institutional policies develop, as scholars accumulate usage experience, and as GenAI tools themselves improve in reliability and transparency.

Fifth, the data-quality issues identified in this dataset junk-encoded demographic rows, backtick-prefixed Likert entries highlight the importance of real-time validation in digital survey platforms. Future surveys should implement constraint rules to prevent encoding artefacts at the point of data entry.

CONCLUSION

This study set out to investigate whether and how information literacy shapes university students' and research scholars' attitudes toward accepting GenAI tools for academic research, within an extended Technology Acceptance Model framework. The answer is clear: Information literacy matters, and it matters independently of and in addition to the perceived usefulness that has long dominated TAM-based explanations of technology adoption.

Drawing on rigorously cleaned survey data from students and research scholars at central and state universities across India, the extended TAM model accounted for roughly two-thirds of the variance in GenAI acceptance attitudes, with information literacy, perceived usefulness, perceived credibility risk, and academic integrity concern all emerging as significant predictors (full statistics are reported in Results). Crucially, demographic characteristics did not differentiate acceptance attitudes, indicating that GenAI

adoption is governed less by who scholars are than by what they know and perceive about these tools. The broader significance of this pattern is that acceptance is a cognitive and perceptual achievement, one that institutions can actively cultivate rather than a fixed attribute of particular groups.

The findings of the study carry a clear message for academic libraries, institutional policymakers, and GenAI tool developers: the path to responsible GenAI adoption in academic research runs through information literacy instruction, transparent communication of research utility, and effective management of credibility risk. As GenAI tools become ever more deeply embedded in the fabric of academic inquiry, investing in information-literate scholars equipped to engage critically, ethically, and strategically with AI-generated content is not merely prudent- it is essential for the integrity and quality of future scholarship.

SIGNIFICANCE STATEMENT

This study extends the Technology Acceptance Model by integrating information literacy, perceived credibility risk, and academic integrity concern to explain the acceptance of generative AI among university students and research scholars in India. The findings demonstrate that information literacy and perceived usefulness significantly promote responsible GenAI acceptance, while credibility risk acts as a barrier to adoption. Unlike demographic characteristics, cognitive competencies and user perceptions primarily drive acceptance attitudes. The study provides practical evidence for academic libraries, higher education institutions, and policymakers to develop information literacy programs and institutional AI guidelines that encourage the ethical and effective use of generative AI in academic research.

REFERENCES

1. Dwivedi, Y.K., N. Kshetri, L. Hughes, E.L. Slade and A. Jeyaraj *et al.*, 2023. Opinion Paper: "So what if ChatGPT wrote it?" Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *Int. J. Inf. Manage.*, Vol. 71. 10.1016/j.ijinfomgt.2023.102642.
2. Kasneci, E., K. Sessler, S. Küchemann, M. Bannert and D. Dementieva *et al.*, 2023. ChatGPT for good? On opportunities and challenges of large language models for education. *Learn. Individual Differ.*, Vol. 103. 10.1016/j.lindif.2023.102274.
3. Davis, F.D., 1989. Perceived usefulness, perceived ease of use and user acceptance of information technology. *MIS Q.*, 13: 319-340.
4. Nazaretsky, T., P. Mejia-Domenzain, V. Swamy, J. Frej and T. Käser, 2025. The critical role of trust in adopting AI-powered educational technology for learning: An instrument for measuring student perceptions. *Comput. Educ.: Artif. Intell.*, Vol. 8. 10.1016/j.caeai.2025.100368.
5. Pillai, R., B. Sivathanu, B. Metri and N. Kaushik, 2024. Students' adoption of AI-based teacher-bots (T-bots) for learning in higher education. *Inf. Technol. People*, 37: 328-355.
6. Archambault, S.G., S. Ramachandran, E. Acosta and S. Fu, 2024. Ethical dimensions of algorithmic literacy for college students: Case studies and cross-disciplinary connections. *J. Acad. Librarianship*, Vol. 50. 10.1016/j.acalib.2024.102865.
7. Venkatesh, V., M.G. Morris, G.B. Davis and F.D. Davis, 2003. User acceptance of information technology: Toward a unified view. *MIS Q.*, 27: 425-478.
8. Eren, B.A., 2024. Antecedents of robo-advisor use intention in private pension investments: An emerging market country example. *J. Financ. Serv. Mark.*, 29: 683-698.
9. Dahri, N.A., N. Yahaya, W.M. Al-Rahmi, A. Aldraiweesh and U. Alturki *et al.*, 2024. Extended TAM based acceptance of AI-Powered ChatGPT for supporting metacognitive self-regulated learning in education: A mixed-methods study. *Heliyon*, Vol. 10. 10.1016/j.heliyon.2024.e29317.
10. Ezeudoka, B.C. and M. Fan, 2024. Determinants of behavioral intentions to use an E-Pharmacy service: Insights from TAM theory and the moderating influence of technological literacy. *Res. Social Administrative Pharm.*, 20: 605-617.

11. Rabiul Awal, M. and M. Enamul Haque, 2025. Revisiting university students' intention to accept AI-Powered chatbot with an integration between TAM and SCT: A South Asian perspective. *J. Appl. Res. Higher Educ.*, 17: 594-608.
12. ALA, 2016. Framework for information literacy for higher education. *Am. Lib. Assoc.*
13. Ray, P.P., 2023. ChatGPT: A comprehensive review on background, applications, key challenges, bias, ethics, limitations and future scope. *Internet Things Cyber-Phys. Syst.*, 3: 121-154.
14. Featherman, M.S. and P.A. Pavlou, 2003. Predicting e-services adoption: A perceived risk facets perspective. *Int. J. Hum. Comput. Stud.*, 59: 451-474.
15. Kaushal, V. and R. Yadav, 2022. The role of chatbots in academic libraries: An experience-based perspective. *J. Aust. Lib. Inf. Assoc.*, 71: 215-232.
16. Wang, X., R. Luo, Y. Liu, P. Chen, Y. Tao and Y. He, 2023. Revealing the complexity of users' intention to adopt healthcare chatbots: A mixed-method analysis of antecedent condition configurations. *Inf. Process. Manage.*, Vol. 60. 10.1016/j.ipm.2023.103444.
17. Stieglitz, S., M. Mirbabaie, A. Deubel, L.M. Braun and T. Kissmer, 2023. The potential of digital nudging to bridge the gap between environmental attitude and behavior in the usage of smart home applications. *Int. J. Inf. Manage.*, Vol. 72. 10.1016/j.ijinfomgt.2023.102665.
18. Dogruel, L., P. Masur and S. Joeckel, 2022. Development and validation of an algorithm literacy scale for internet users. *Commun. Methods Meas.*, 16: 115-133.
19. Trepte, S., D. Teutsch, P.K. Masur, C. Eicher, M. Fischer, A. Hennhöfer and F. Lind, 2014. Do People Know About Privacy and Data Protection Strategies? Towards the "Online Privacy Literacy Scale" (OPLIS). In: *Reforming European Data Protection Law*, Gutwirth, S., R. Leenes and P. de Hert (Eds.), Springer, Netherlands, ISBN: 978-94-017-9385-8, pp: 333-365.
20. Nikou, S., M. de Reuver and M.M. Kanafi, 2022. Workplace literacy skills-how information and digital literacy affect adoption of digital technology. *J. Doc.*, 78: 371-391.
21. Cetindamar, D., B. Abedin and K. Shirahada, 2021. The role of employees in digital transformation: A preliminary study on how employees' digital literacy impacts use of digital technologies. *IEEE Trans. Eng. Manage.*, 71: 7837-7848.
22. Shin, D., A. Rasul and A. Fotiadis, 2022. Why am I seeing this? Deconstructing algorithm literacy through the lens of users. *Internet Res.*, 32: 1214-1234.
23. Safadel, P., S.N. Hwang and J.M. Perrin, 2023. User acceptance of a virtual librarian chatbot: An implementation method using IBM Watson natural language processing in virtual immersive environment. *TechTrends*, 67: 891-902.
24. DeLone, W.D. and E.R. McLean, 2003. The DeLone and McLean model of information systems success: A ten-year update. *J. Manage. Inf. Syst.*, 19: 9-30.
25. Hair, J.F., W.C. Black, B.J. Babin and R.E. Anderson, 2019. *Multivariate Data Analysis*. 8th Edn., Cengage, Boston, USA, ISBN: 9781473756540, Pages: 813.
26. Noguera-Vivo, J.M., M. del Mar Grandío-Pérez, 2025. Enhancing algorithmic literacy: Experimental study on communication students' awareness of algorithm-driven news. *Anàlisi*, 71: 37-53.