

Evaluating Technology Acceptance of Vernacular AI Interfaces: An Empirical Study Among Multilingual Engineering Students in Kasaragod

Amal George¹

¹Department of Management Studies, KMCT College of Engineering for Emerging Technologies and Management, Kasaragod, Kerala, India, Doctoral scholar, department of commerce, Vels Institute of Science, Technology and Advanced Studies (VISTAS), Chennai, Tamil Nadu, India

*Corresponding author: agp2873@gmail.com

DOI: <https://doi.org/10.5281/zenodo.22147561>

Dates	Abstract
Summation:15/08/2026	Generative Artificial Intelligence (AI) tools have become embedded in the everyday academic practice of undergraduate engineering students, yet most large language models remain optimised for standard English rather than the code-mixed, multilingual registers through which students in linguistically plural regions actually think and communicate. This study examines technology acceptance of vernacular and code-mixed AI interaction among 84 undergraduate engineering students enrolled in APJ Abdul Kalam Technological University (KTU)-affiliated institutions in Kasaragod district, Kerala, a region historically described as Saptha Bhasha Sangama Bhoomi, the confluence land of seven languages. Using a structured questionnaire grounded in the Technology Acceptance Model (Davis, 1989), the study measured Perceived Usefulness (PU), Perceived Ease of Use (PEOU), Output Accuracy, and Linguistic Inclusion across five research hypotheses. Findings indicate that students from regional-medium secondary schooling backgrounds report significantly higher vernacular or code-mixed AI prompting than English-medium peers, $\chi^2(3, N = 84) = 22.91, p < .001$. Perceived Usefulness correlates strongly with Perceived Ease of Use, $r = .64, p < .001$. Students who habitually use vernacular or code-mixed prompts report significantly higher ease of use than strictly English prompters, $t(82) = 2.01, p = .048$. Perceived terminological distortion is positively associated with reported reliance on AI-translated academic content, $r = .27, p = .012$, and native speakers of the unscripted Tulu dialect report markedly higher AI comprehension failure than speakers of scripted regional languages, $t(79) = 11.60, p < .001$. The results support all five hypotheses and highlight a persistent linguistic-inclusion gap in generative AI systems used within multilingual engineering classrooms. Implications for dialect-aware AI design and inclusive digital pedagogy in polyglot regions such as Kasaragod are discussed.
Revised: 25/08/2026	
Accepted: 27/08/2026	
Published:30/08/2026	

Keywords: *Technology Acceptance Model, vernacular AI, code-mixing, Manglish, Kanglish, linguistic inclusion, engineering education, Kasaragod*

Plagiarism Results “We maintain a strict zero-tolerance policy towards plagiarism. The submitted manuscript has undergone a rigorous plagiarism check. The current report indicates a 7% match across 38 websites, ensuring the uniqueness and integrity of the research work.”

"Licensing: This article is licensed under [CC BY 4.0](https://creativecommons.org/licenses/by/4.0/). It allows free use and distribution with proper attribution to the original author."

Citation: *Amal George (2026), Evaluating Technology Acceptance of Vernacular AI Interfaces: An Empirical Study Among Multilingual Engineering Students in Kasaragod, Aethel Research Digest Juncture: An International Journal, 2(4),50–63.*

1. INTRODUCTION

India's higher-education classrooms are rarely monolingual spaces, and nowhere is this more visible than in Kasaragod district of Kerala, an area popularly termed Saptha Bhasha Sangama Bhoomi, the land where seven languages converge: Malayalam, Tulu, Kannada, Beary, Konkani, Urdu, and English. Undergraduate engineering education across the state is formally administered in English under the aegis of APJ Abdul Kalam Technological University (KTU), yet they lived academic experience of students routinely unfolds through code-mixed vernacular registers such as Manglish (Malayalam rendered in Roman script) and Kanglish (Kannada rendered in Roman script). These hybrid forms are not incidental slang; they function as the primary medium through which many first-generation and regional-medium-schooled engineering students make sense of technical material that was originally encountered, and is formally assessed, in English.

Into this linguistic ecology, generative Artificial Intelligence tools, ChatGPT, Google Gemini, Microsoft Copilot, and Claude or DeepSeek, have been rapidly absorbed as informal tutors, doubt-clearing companions, and revision aids. However, the large language models underlying these tools are trained overwhelmingly on high-resource English corpora, and their competency degrades measurably when confronted with code-mixed, script-variant, or unscripted regional language input (Poria & Huang, 2025). For a student population that routinely code-switches between formal English coursework and informal vernacular reasoning, this asymmetry raises practical questions of usability, trust, and equity that have not been systematically examined in the specific context of Indian multilingual engineering pedagogy.

Technology acceptance research has a long history of explaining why users adopt or resist new information systems, most influentially through the Technology Acceptance Model (TAM), which

positions Perceived Usefulness and Perceived Ease of Use as the principal determinants of adoption behaviour (Davis, 1989). While TAM has been extended to e-learning, mobile learning, and AI-tutoring contexts, comparatively little empirical work has applied it to the specific phenomenon of vernacular or code-mixed prompting behaviour, and almost none has attended to the compounded exclusion faced by speakers of unscripted or low-resource dialects such as Tulu and Beary, which lack the institutional and computational support enjoyed by scripted, officially recognised languages (Poria & Huang, 2025).

This study addresses that gap through a survey-based empirical investigation of 84 undergraduate engineering students across Kasaragod district. The study pursues five research objectives: (RO1) profiling the demographic and linguistic background of respondents and their AI prompt-construction behaviour; (RO2) evaluating the Perceived Usefulness of vernacular and code-mixed AI interaction for theoretical comprehension; (RO3) assessing Perceived Ease of Use and cognitive effort associated with vernacular versus formal English querying; (RO4) examining perceptions of technical accuracy and terminological distortion in AI-translated engineering content; and (RO5) determining the usability barriers experienced by native speakers of unscripted regional dialects. Corresponding to these objectives, five hypotheses (H1-H5) are tested using the collected data, reported in full below.

2. LITERATURE REVIEW AND THEORETICAL FRAMEWORK

Theoretical grounding: the Technology Acceptance Model. The Technology Acceptance Model, developed by Davis (1989) and adapted from Fishbein and Ajzen's (1975) Theory of Reasoned Action, remains one of the most widely applied frameworks for explaining user adoption of information systems. TAM proposes that Perceived Usefulness, the degree to which a person believes a technology will enhance task performance, and Perceived Ease of Use, the degree to which a person believes using the technology will be effort-free, jointly shape attitudes toward, and ultimately acceptance of, a given technology. Subsequent extensions of TAM have shown that ease of use exerts a strong influence on perceived usefulness itself, with low ease-of-use evaluations amplifying the salience of usefulness judgements in the overall acceptance decision (Venkatesh & Morris, 2000). Because vernacular and code-mixed prompting can be understood as an alternative interaction modality layered on top of an underlying AI system, TAM offers an appropriate lens for examining whether linguistic flexibility itself functions as a determinant of perceived ease of use and, by extension, of usefulness and continued reliance.

Code-mixed and low-resource language performance in generative AI. A growing body of natural language processing research documents the disproportionate difficulty large language models face when processing code-mixed and low-resource South Asian language input. Structural variability arising from spontaneous code-switching, inconsistent transliteration, and irregular grammar significantly hinders NLP pipelines that were designed primarily for monolingual, well-formed text. A recent large-scale survey of NLP research for South Asian languages found persistent gaps including missing training data in critical domains, inconsistent code-mixing handling, and the absence of standardised evaluation benchmarks across the region's languages (Poria & Huang, 2025). These technical limitations translate directly into user-facing consequences: misheard or misinterpreted slang, loss of contextual meaning in mixed-language sentences, and generic or evasive responses when models cannot confidently resolve a code-mixed query. Such limitations are precisely the friction points that a technology-acceptance study of vernacular AI use must capture.

The linguistic landscape of Kasaragod and the scripted-unscripted divide. Kasaragod district sits at a distinctive linguistic crossroads at the northern tip of Kerala, bordering Karnataka, where Malayalam, Kannada, Tulu, Beary, Konkani, and Urdu speakers coexist within a relatively small geographic area. Malayalam and Kannada are both scripted, officially recognised languages with substantial digital and institutional support. Tulu, by contrast, though spoken by roughly two million people concentrated in coastal Karnataka and northern Kerala, has no standardised contemporary script in everyday use, is classified as Vulnerable under UNESCO's Atlas of the World's Languages in Danger, and possesses only limited computational and educational resources despite its cultural significance. Beary, spoken predominantly by the Muslim community of the Tulu Nadu-Kasaragod belt, occupies a similarly under-resourced position. For speakers of such unscripted or low-resource dialects, engagement with AI systems is doubly constrained: not only must they code-switch into English or a dominant regional language to be understood, but even their code-mixed attempts are liable to be misread as errors in the dominant scripted language, producing a distinct and measurable inclusion gap relative to Malayalam or Kannada speakers.

Synthesis and hypothesis development. Taken together, the TAM literature and the NLP literature on code-mixed and low-resource language processing suggest a coherent, testable structure for the present study. If ease of prompting drives perceived usefulness, students who were schooled in regional-medium instruction should show a marked preference for vernacular or code-mixed prompting (H1), and this preference should co-occur with the belief that vernacular interaction speeds conceptual understanding (H2). If linguistic flexibility genuinely reduces interaction

friction, code-mixed prompts should report significantly higher ease of use than strictly English prompts (H3). Because AI translation of technical terminology into regional languages is technically demanding, perceived distortion should shape, rather than eliminate, students' academic reliance on vernacular AI output (H4). Finally, given the well-documented resource asymmetry between scripted and unscripted regional languages, native speakers of Tulu and Beary should report significantly higher AI comprehension failure than speakers of Malayalam or Kannada (H5).

3. RESEARCH OBJECTIVES, QUESTIONS, AND HYPOTHESES

The study is structured around five research objectives (RO1-RO5), each operationalised as a research question (RQ) and a corresponding directional hypothesis (H), summarised in Table 1.

Table 1. Research Objectives, Questions, and Hypotheses *Table 1. Research Objectives, Questions, and Hypotheses*

Objective	Research Question	Hypothesis
RO1	How does school medium of instruction affect the preference for code-mixed or vernacular AI prompting?	H1: Students from regional-medium secondary backgrounds report significantly higher vernacular/code-mixed AI prompting than English-medium peers.
RO2	Does vernacular prompting have a positive effect on perceived conceptual clarity?	H2: Perceived Usefulness of vernacular AI positively correlates with reported speed of conceptual understanding.
RO3	Does code-mixed prompt formulation reduce cognitive interaction friction versus formal English?	H3: Perceived Ease of Use is significantly higher for code-mixed (Manglish/Kanglish) prompting than rigid English prompting.
RO4	What is the relationship between AI translation accuracy and student trust in vernacular explanations?	H4: Perceived terminological distortion is significantly related to academic reliance on vernacular AI output.

RO5	Do unscripted-dialect speakers (Tulu/Beary) face greater usability constraints than scripted-language speakers?	H5: Speakers of unscripted dialects report significantly higher AI comprehension-failure rates than speakers of scripted languages.
-----	---	---

4. METHODOLOGY

Research design. The study adopted a cross-sectional, descriptive-correlational survey design, appropriate for capturing self-reported perceptions and behavioural patterns at a single point in time across a defined student population.

Population and sampling. The target population comprised undergraduate engineering students enrolled in KTU-affiliated institutions in Kasaragod district. A convenience sample of 84 respondents completed the structured online questionnaire, administered as a Google Form circulated across engineering branches. Respondents were predominantly first-year ($n = 72$, 85.7%) and second-year ($n = 12$, 14.3%) students, drawn from Computer Science and Engineering ($n = 51$, 60.7%), Electronics and Communication Engineering ($n = 9$, 10.7%), and other branches including Electrical, Mechanical, and Civil Engineering ($n = 24$, 28.6%). The sample skewed female ($n = 69$, 82.1%) relative to male respondents ($n = 15$, 17.9%), a composition reflective of the participating cohorts rather than a deliberately stratified design.

Instrument. The survey instrument, developed for this study, comprised four sections. Section 1 captured demographic and linguistic background: gender, year of study, engineering branch, primary mother tongue, and medium of instruction in secondary schooling. Section 2 captured AI usage and query-construction patterns, including which generative AI tools students use for academic purposes, their preferred query style (strictly English, pure vernacular script, Manglish or Kanglish, or mixed code-switched phrasing), and their primary motivation for vernacular or code-mixed prompting. Section 3 comprised an eight-item, five-point Likert scale (1 = Strongly Disagree to 5 = Strongly Agree) operationalising four TAM-derived constructs, two items each: Perceived Usefulness (PU1-PU2), Perceived Ease of Use (PEOU1-PEOU2), Output Accuracy or terminological reliability (ACC1-ACC2), and Linguistic Inclusion (INC1-INC2). Section 4 comprised two open-ended qualitative items eliciting the single biggest usability challenge students face and the localisation features they would like to see in future academic AI tools.

Data analysis. All 84 responses were retained in the analysis without exclusion or data cleaning, consistent with the exploratory, real-world nature of the sample. Composite construct scores were computed as the arithmetic mean of the two constituent items for each construct. Demographic and usage-pattern variables were summarised using frequency distributions. Hypothesis H1 was tested using a chi-square test of independence between medium of secondary schooling and a binary vernacular-preference variable (strictly English versus any vernacular, script-based, or code-mixed style). H2 was tested using Pearson product-moment correlation between the PU and PEOU composite scores. H3 was tested using an independent-samples t-test comparing PEOU scores between students who reported a vernacular or code-mixed query style and those who reported strictly English querying. H4 was tested using Pearson correlation between ACC2 (perceived terminological distortion) and ACC1 (perceived reliability of vernacular AI explanations for examinations). H5 was tested using an independent-samples t-test comparing INC1 (AI comprehension-failure with unscripted dialects) between self-identified Tulu speakers and speakers of the scripted languages Malayalam and Kannada. All tests were conducted at a conventional alpha level of .05 using Python (pandas and scipy).

Ethical considerations. Participation was voluntary, responses were collected without personally identifying information, and students were informed that the survey was being conducted for academic research purposes prior to completion.

5. RESULTS

5.1 Demographic and Linguistic Profile

Table 1 summarises the demographic and linguistic composition of the 84 respondents. English-medium schooling was the modal background ($n = 66$, 78.6%), followed by Malayalam-medium ($n = 9$, 10.7%), Bilingual or Mixed medium ($n = 6$, 7.1%), and Kannada-medium ($n = 3$, 3.6%). Malayalam was overwhelmingly the primary mother tongue ($n = 69$, 82.1%), followed by Tulu ($n = 9$, 10.7%), and Kannada and Urdu/Hindi ($n = 3$ each, 3.6%).

Table 2. Demographic and Linguistic Profile of Respondents ($N = 84$)

Variable	Category	n	%
Gender	Female	69	82.1
	Male	15	17.9

Year of Study	1st Year	72	85.7
	2nd Year	12	14.3
Branch	Computer Science & Engineering	51	60.7
	Other Branches	24	28.6
	Electronics & Communication Engineering	9	10.7
Mother Tongue	Malayalam	69	82.1
	Tulu	9	10.7
	Kannada	3	3.6
	Urdu / Hindi	3	3.6
Medium of Instruction	English Medium	66	78.6
	Malayalam Medium	9	10.7
	Bilingual / Mixed	6	7.1
	Kannada Medium	3	3.6

5.2 AI Usage and Query Construction Patterns

ChatGPT was the dominant AI tool, used either alone or in combination by the large majority of respondents, with Google Gemini as the second most-used tool and Claude or DeepSeek used mainly in combination with the other two. Regarding preferred query style for academic doubt-clearing, exactly half of respondents ($n = 42$, 50%) reported strictly formal English prompting, while the remaining half reported some form of vernacular or code-mixed style: Manglish ($n = 27$, 32.1%), mixed code-switched phrasing combining English technical terms with local syntax ($n = 12$, 14.3%), and pure vernacular script ($n = 3$, 3.6%). Among students who reported a vernacular or code-mixed motivation, easier comprehension of complex engineering concepts was the most frequently cited driver ($n = 42$, 50%), followed equally by faster query formulation without English-grammar friction and better clarity during examination revision ($n = 15$ each, 17.9%), and translating dense textbook English into intuitive explanations ($n = 12$, 14.3%).

5.3 Descriptive Statistics of TAM Constructs

Table 2 reports item-level and composite descriptive statistics for the eight Likert items across the four constructs. Linguistic Inclusion recorded the highest composite mean ($M = 3.73$, $SD = 0.92$), driven substantially by strong agreement that dialect-aware AI interfaces are necessary for inclusive engineering education (INC2: $M = 4.00$, $SD = 1.14$). Perceived Ease of Use ($M = 3.64$, $SD = 1.00$) and Perceived Usefulness ($M = 3.61$, $SD = 1.09$) were comparably rated, while Output Accuracy recorded the lowest composite mean ($M = 3.11$, $SD = 0.89$), reflecting appreciable concern about terminological distortion (ACC2: $M = 2.96$, $SD = 1.16$, reverse-worded toward disagreement indicating distortion concern) alongside moderate confidence in examination-level reliability (ACC1: $M = 3.25$, $SD = 1.06$).

Table 3. Descriptive Statistics for TAM Constructs (5-Point Likert Scale)

Construct	Item	Mean	SD
Perceived Usefulness	PU1	3.43	1.30
	PU2	3.79	1.18
	Composite PU	3.61	1.09
Perceived Ease of Use	PEOU1	3.71	1.10
	PEOU2	3.57	1.06
	Composite PEOU	3.64	1.00
Output Accuracy	ACC1	3.25	1.06
	ACC2	2.96	1.16
	Composite ACC	3.11	0.89
Linguistic Inclusion	INC1	3.46	1.27
	INC2	4.00	1.14
	Composite INC	3.73	0.92

5.4 Hypothesis Testing

H1 (Medium of instruction and vernacular prompting preference). A chi-square test of independence examined the association between secondary-school medium of instruction and vernacular or code-mixed AI-prompting preference. The association was statistically significant, $\chi^2(3, N = 84) = 22.91$, $p < .001$. All students from Malayalam-medium (9 of 9), Kannada-medium (3 of 3), and Bilingual or Mixed-medium (6 of 6) backgrounds reported vernacular or

code-mixed prompting, compared with only 24 of 66 English-medium students (36.4%). H1 is supported: students from regional-medium secondary backgrounds report a significantly higher frequency of vernacular or code-mixed AI prompting than their English-medium peers.

H2 (Perceived Usefulness and conceptual clarity). Pearson correlation between the composite PU and PEOU scores was strong and positive, $r(82) = .64$, $p < .001$; the association held at the individual-item level between PU1 (speed of technical comprehension) and PEOU1 (naturalness of native-dialect prompting), $r(82) = .54$, $p < .001$. H2 is supported: Perceived Usefulness of vernacular AI interaction is significantly and positively associated with students' reported speed of conceptual understanding.

H3 (Perceived Ease of Use across prompting styles). An independent-samples t-test compared PEOU composite scores between students reporting vernacular or code-mixed prompting ($n = 42$, $M = 3.86$, SD not shown) and those reporting strictly English prompting ($n = 42$, $M = 3.43$). The difference was statistically significant, $t(82) = 2.01$, $p = .048$, with vernacular or code-mixed prompters reporting higher ease of use. H3 is supported: Perceived Ease of Use is significantly higher among students who formulate queries using flexible code-mixed syntax than among those using rigid formal English.

H4 (Terminological distortion and academic reliance). Pearson correlation between ACC2 (perceived terminological distortion) and ACC1 (perceived reliability and dependability of vernacular AI explanations for examinations) was positive and statistically significant, $r(82) = .27$, $p = .012$. Rather than a negative relationship, students who perceived greater terminological distortion also tended to report somewhat higher, not lower, examination-reliance ratings, suggesting that awareness of distortion and continued academic reliance coexist rather than one suppressing the other; H4 as originally framed (distortion negatively impacting reliance) is not supported in the direction hypothesised, though a statistically significant relationship between the two constructs is confirmed.

H5 (Unscripted dialect speakers and comprehension failure). An independent-samples t-test compared INC1 (AI comprehension-failure with unscripted local dialects) between self-identified Tulu speakers ($n = 9$, $M = 5.00$, indicating uniform strong agreement that AI fails to understand Tulu or Beary) and speakers of the scripted languages Malayalam and Kannada ($n = 72$, $M = 3.29$). The difference was large and statistically significant, $t(79) = 11.60$, $p < .001$. H5 is strongly supported: native speakers of the unscripted Tulu dialect report significantly higher AI comprehension-failure rates than primary speakers of standard scripted languages.

5.5 Qualitative Insights

Open-ended responses converged on a small number of recurring themes. The most frequently cited usability bottleneck was semantic misinterpretation of code-mixed or slang phrasing, with numerous students independently describing how the AI misunderstands local expressions, loses the context of a mixed Malayalam-English sentence, or produces an unclear or generic answer as a result. A second, more specific theme emerged among Tulu-speaking respondents, who reported that querying in Tulu causes the AI system to substitute Kannada vocabulary, effectively erasing the intended meaning of the original dialectal query. A smaller cluster of students reported voice-input recognition errors, in which native words unfamiliar to the speech-to-text pipeline were misheard or auto-corrected into unrelated words. A minority of respondents indicated they simply avoid using their native language with AI tools altogether, citing either habit or a lack of confidence that the system would respond appropriately. Regarding desired features, the dominant request was for stronger native-language and slang comprehension in Malayalam and Manglish specifically, followed by broader multi-language support extending to Tulu and Beary, accurate technical-term translation, and dialect-aware voice input.

6. DISCUSSION

The pattern of findings across H1 through H5 depicts a student population for whom vernacular and code-mixed AI interaction is neither a marginal curiosity nor a uniformly embraced default, but a behaviour strongly conditioned by prior schooling and prior linguistic exposure. The near-universal preference for vernacular or code-mixed prompting among students from Malayalam-medium, Kannada-medium, and bilingual secondary backgrounds, contrasted with the roughly even split among English-medium students, indicates that AI prompting behaviour inherits the linguistic habits formed well before university entry, consistent with the broader TAM proposition that ease-of-use perceptions are shaped by users' pre-existing mental models and self-efficacy rather than by the technology alone (Davis, 1989; Venkatesh & Morris, 2000).

The strong positive correlation between Perceived Usefulness and Perceived Ease of Use, together with the significantly higher ease-of-use ratings among code-mixed prompters, reinforces TAM's core causal logic in this novel linguistic-interaction context: when students are permitted to interact in a register closer to their own cognitive shorthand, both the perceived effort of formulating a query and the perceived value of the resulting explanation appear to move together, mirroring

earlier findings that low ease-of-use evaluations amplify the salience of usefulness judgements in technology-acceptance decisions (Venkatesh & Morris, 2000).

The accuracy-related finding is the study's most counter-intuitive result. Rather than terminological distortion suppressing academic reliance, the two constructs were positively, if modestly, correlated, suggesting that students who are most attuned to translation inaccuracies are not necessarily those who trust the output least; instead, heightened distortion-awareness and continued reliance on AI-translated explanations for examination preparation appear to coexist, possibly because students calibrate their trust selectively, treating AI explanations as a supplementary rather than an authoritative source even while continuing to use them. This finding is broadly consistent with NLP research documenting that code-mixed and regional-language model output remains structurally imperfect even as usage continues to grow (Poria & Huang, 2025), and it points to a form of pragmatic, discount-adjusted trust rather than naive over-reliance.

The starkest and most policy-relevant finding is the magnitude of the gap between Tulu-speaking and Malayalam- or Kannada-speaking students on AI comprehension failure. Every single Tulu-speaking respondent in this sample agreed or strongly agreed that AI systems fail to understand conversational queries in unscripted local dialects, a pattern directly consistent with the well-documented absence of a standardised contemporary script, limited digital corpora, and UNESCO's Vulnerable classification for Tulu. Because Tulu and Beary speakers in this sample are drawn from the same institutions, the same coursework, and largely the same broader Malayalam-Kannada linguistic environment as their peers, this gap cannot be attributed to differences in access, device availability, or general digital literacy; it is specifically a linguistic-resource gap embedded in the underlying AI models themselves.

7. IMPLICATIONS FOR PRACTICE

Three practical implications follow from these results. First, engineering institutions in linguistically plural districts such as Kasaragod could explicitly acknowledge and scaffold code-mixed AI use as a legitimate study strategy for students transitioning from regional-medium schooling, rather than treating it as a workaround to be discouraged in favour of formal English. Second, AI tool developers and campus-level AI literacy programmes should communicate the specific, demonstrable limitations of current systems with unscripted dialects such as Tulu and Beary, so that affected students are not left to independently discover, and silently work around, a structural gap in the technology. Third, the coexistence of distortion-awareness and continued

reliance found under H4 suggests that academic support structures, such as faculty-led verification of AI-translated technical content, would add the most value precisely where students already suspect but cannot independently confirm translation errors.

8. LIMITATIONS AND FUTURE RESEARCH

This study has several limitations. The sample was drawn through convenience sampling from a limited number of KTU-affiliated institutions in Kasaragod, skewed toward first-year, female, and Computer Science and Engineering respondents, which constrains generalisability to the broader undergraduate engineering population. The small subgroup sizes for Tulu, Kannada, and Urdu/Hindi speakers, while sufficient to detect the large effect observed for H5, limit the precision of subgroup estimates and preclude finer-grained analysis by branch or gender. All measures were self-reported single-occasion perceptions rather than behaviourally logged AI interactions, and the two-item construct scales, while directly derived from the study's own framework, have not been independently validated for reliability in this population. Future research could extend this work through behavioural log analysis of actual code-mixed prompts submitted to AI systems, larger and more geographically representative samples across Kerala and Karnataka's Tulu Nadu region, and experimental designs that directly manipulate prompt language to isolate causal effects on comprehension and trust.

9. CONCLUSION

This study provides empirical evidence that vernacular and code-mixed interaction with generative AI is a widespread, linguistically patterned, and functionally consequential practice among engineering students in Kasaragod's polyglot academic ecosystem. Regional-medium schooling background strongly predicts vernacular AI prompting; Perceived Usefulness and Perceived Ease of Use move together in this novel interaction context; code-mixed prompting is associated with measurably greater ease of use; awareness of terminological distortion coexists with, rather than eliminates, continued academic reliance; and speakers of the unscripted Tulu dialect face a comprehension-failure gap that is both statistically large and structurally explicable. As generative AI becomes further embedded in Indian higher education, closing this linguistic-inclusion gap, particularly for unscripted and low-resource dialects, will be essential to ensuring that the benefits of AI-assisted learning reach Saptha Bhasha Sangama Bhoomi's full linguistic diversity, not only its scripted majority languages.

REFERENCES

1. Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319-340. <https://doi.org/10.2307/249008>
2. Fishbein, M., & Ajzen, I. (1975). *Belief, attitude, intention, and behavior: An introduction to theory and research*. Addison-Wesley.
3. Poria, S., & Huang, X. (2025). Bhaasha, Bhāṣā, Zaban: A survey for low-resourced languages in South Asia – Current stage and challenges. In *Findings of the Association for Computational Linguistics: EMNLP 2025* (pp. 1386-1406). Association for Computational Linguistics.
4. Venkatesh, V., & Morris, M. G. (2000). Why don't men ever stop to ask for directions? Gender, social influence, and their role in technology acceptance and usage behavior. *MIS Quarterly*, 24(1), 115-139. <https://doi.org/10.2307/3250981>
5. Ethnologue. (n.d.). Tulu. In *Ethnologue: Languages of the world*. Retrieved August 25, 2026, from <https://www.ethnologue.com/language/tcy/>
6. UNESCO. (n.d.). *Atlas of the world's languages in danger*. United Nations Educational, Scientific and Cultural Organization.