
| RESEARCH ARTICLE

Algorithmically Mediated Audience Design: Voice, Pragmatic Force, and Linguistic Normativity in Generative-AI Rewriting

Thamer Marzouq¹, Naif Alqurashi²

¹ Assistant Professor, Department of Foreign Languages, Taif University, Taif, Kingdom of Saudi Arabia

² Associate Professor, Department of Foreign Languages, Taif University, Taif, Kingdom of Saudi Arabia

Corresponding Author: Naif Alqurashi; **E-mail:** dr.naif.alqurashi@tu.edu.sa

| ABSTRACT

Generative AI is increasingly inserted between multilingual writers and their audiences, yet its linguistic consequences are often discussed as matters of accuracy, fluency, or efficiency. This study reconceptualizes AI-assisted rewriting as algorithmically mediated audience design: a process in which a generative system can redistribute pragmatic force, epistemic commitment, register, authorial visibility, and the indexical meanings through which writers position themselves. A controlled qualitative elicitation corpus was constructed from 12 researcher-designed English utterances representing recurrent applied-linguistic phenomena, including directness, stance, culturally situated formulaicity, pluricentric lexis, code-meshing, self-mention, and relational deference. Each item was contrasted across generic “improvement” and audience/identity-aware revision conditions. Analysis combined functional linguistic coding with deviant-case analysis and an explicit audit trail across lexicogrammatical, pragmatic, discourse, and sociolinguistic levels. Five recurrent processes were identified: sociolinguistic normalization, pragmatic mitigation, epistemic recalibration, redistribution of authorial voice, and selective retention of multilingual indexicality. Audience-explicit prompting constrained some homogenizing tendencies but did not eliminate normative mediation. The findings support a model of algorithmically mediated audience design in which prompting, model output, and human uptake jointly shape communicative meaning. The study argues that applied-linguistic evaluation of GenAI should move beyond error reduction toward consequential changes in voice, identity, normativity, and communicative agency.

| KEYWORDS

Generative AI; audience design; sociolinguistics; pragmatics; multilingual writing; authorial voice; linguistic normativity; applied linguistics

| ARTICLE INFORMATION

ACCEPTED: 06 August 2026

PUBLISHED: 16 September 2026

DOI: 10.32996/jpda.2026.5.5.5

1. Introduction

Generative artificial intelligence has rapidly become part of ordinary language work. Multilingual writers use large language models to reformulate requests, translate culturally situated expressions, adjust academic prose, soften disagreement, generate alternatives, and anticipate how a text may be received. These practices appear superficially similar to proofreading or editing, but they are linguistically more consequential. A revision that replaces a direct request with a modal interrogative changes interpersonal force; a revision that removes first-person self-mention changes authorial positioning; a revision that substitutes a pluricentric expression with a dominant standardized form changes the sociolinguistic repertoire presented to the audience. The relevant object for applied linguistics is therefore not merely whether AI produces “better English,” but how AI-mediated rewriting reorganizes meaning, social relations, and legitimate ways of sounding like an English user.

This problem sits at the intersection of several established applied-linguistic traditions. Bell’s audience design treats stylistic variation as responsive to addressees and other audience roles. Pragmatics examines how linguistic choices enact entitlement, politeness, stance, and social distance. Research on academic discourse demonstrates that hedges, boosters, self-mention, and evaluative resources are constitutive of disciplinary voice rather than decorative style. Translanguaging and Global Englishes

Copyright: © 2026 the Author(s). This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC-BY) 4.0 license (<https://creativecommons.org/licenses/by/4.0/>). Published by Al-Kindi Centre for Research and Development, London, United Kingdom.

scholarship, meanwhile, challenges the assumption that communicative legitimacy is exhausted by conformity to a single standardized norm. GenAI brings these traditions into a new configuration because an algorithmic intermediary can modify a text before the human audience encounters it.

Recent work in Applied Linguistics makes the timing of this question particularly salient. The journal's permanent AI in Applied Linguistics section explicitly calls for research in which AI illuminates durable questions about meaning-making, interaction, multilingual practice, communicative competence, and epistemic access rather than functioning as a transient product demonstration (Martinez, Martinez, & Curle, 2026). Tang's (2026) concept of AI-textuality theorizes human-AI textual production as an expanded intertextual relation; Khuder (2026) shows that disciplinary voice in AI-assisted writing depends on critical feedback-seeking; and Wang et al. (2026) demonstrate that advanced multilingual writers position AI variably as tool, interlocutor, and distributed agent. At the same time, emerging work on digital linguistic imperialism warns that AI writing infrastructures may pathologize or marginalize non-dominant linguistic resources (Mohammad Hosseinpour, 2026). Together, these developments suggest that AI-mediated revision should be examined as a sociolinguistic and pragmatic process, not only as educational technology.

The present study addresses this problem through a controlled qualitative elicitation corpus. Its purpose is deliberately bounded. It does not claim to measure learner development, user attitudes, or classroom effectiveness, because the dataset contains no human participants. Instead, it asks what happens linguistically when marked L2, pluricentric, culturally situated, and interpersonal expressions are submitted to generic versus audience-explicit rewriting instructions. This design makes it possible to isolate transformation relations between source and output while avoiding unsupported claims about cognition. The study advances the construct of algorithmically mediated audience design (AMAD), defined as the redistribution of linguistic choices through interaction among a writer's projected audience, prompt design, model-generated alternatives, and subsequent human uptake.

Three research questions guide the analysis: (RQ1) What recurrent linguistic transformations occur when GenAI is asked to improve L2, pluricentric, or pragmatically marked English for academic/professional use? (RQ2) How do these transformations affect pragmatic force, epistemic stance, authorial voice, and sociolinguistic identity? (RQ3) How does explicit instruction to preserve audience relationship, cultural positioning, or multilingual voice alter the model's revisions?

2. Literature Review

2.1 Audience design as an applied-linguistic starting point

Audience design offers a productive starting point because it conceptualizes style as socially responsive. Bell (1984, 2001) argues that speakers shift linguistic style in relation to addressees and, to varying degrees, auditors, overhearers, and referees. The importance of this model for digital writing lies in its refusal to treat linguistic form as autonomous from projected reception. Writers select forms partly in anticipation of who will interpret them and what social relationship the text should instantiate. Digital communication intensifies this anticipatory dimension because audiences may be multiple, collapsed, partly imagined, and dynamically assembled.

GenAI adds an intermediary audience-like orientation without becoming an audience in the human sociological sense. Writers learn that particular prompts elicit particular registers, levels of directness, or forms of standardization; they therefore design instructions for the model while simultaneously designing the eventual text for a human recipient. The resulting text is doubly oriented. It is produced through interaction with a computational system whose anticipated behavior matters to the writer, and it is subsequently evaluated by human addressees. AMAD extends audience design by making this intermediate transformation analytically visible.

2.2 Pragmatic force, stance, and authorial voice

AI revision can alter propositional meaning indirectly by modifying illocutionary force and epistemic commitment. Requests vary in entitlement and mitigation; disagreements can be strengthened or softened; evidential claims can be recast through modal verbs, reporting verbs, adverbs, and hedges. In academic discourse, such resources are central to the negotiation of knowledge and responsibility. Replacing "I strongly believe" with an impersonal declarative, for example, is not simply stylistic cleaning: it changes who is textually accountable for the claim. Likewise, replacing "definitely" with "may" changes the degree of commitment encoded in the proposition.

This matters for multilingual writers because linguistic advice is often presented through the apparently neutral vocabulary of clarity, professionalism, and naturalness. Those labels can obscure normative judgments about how much warmth, deference, directness, self-mention, or cultural markedness is acceptable. Khuder's (2026) work on disciplinary voice is instructive here: AI assistance becomes pedagogically valuable when writers actively regulate feedback rather than treating generated revisions as

authoritative. The present study extends that insight from disciplinary writing to a wider set of pragmatic and sociolinguistic resources.

2.3 Multilingualism, normativity, and indexical meaning

From a multilingual perspective, “improvement” is an ideologically loaded instruction. English is pluricentric and socially stratified; forms that are marked in one context may be conventional in another. Translanguaging further emphasizes that multilingual speakers mobilize resources across named-language boundaries to accomplish identity and interactional work. When an AI system deletes an Arabic discourse marker, replaces a South Asian professional expression, or compresses a culturally elaborate relational formula, it may increase conformity to a dominant professional norm while simultaneously reducing indexical information about affiliation, stance, or cultural positioning.

The concept of homogeneity-by-design is useful because it shifts attention from isolated errors to patterned reduction of variation. Alvero et al. (2025) argue that digital writing infrastructures can compress stylistic diversity, while Mohammad Hosseinpur (2026) provides recent evidence that AI writing tools can reproduce linguistic hierarchies in EFL contexts. The theoretical issue is not that every standardized revision is undesirable. Accommodation is often communicatively strategic. The issue is whether writers can recognize what is being normalized and retain the agency to accept, reject, gloss, or strategically preserve marked resources.

2.4 A model of algorithmically mediated audience design

AMAD conceptualizes AI-assisted rewriting as a four-stage recursive process. First, the writer projects a human audience and communicative purpose. Second, that projection is translated into a prompt, either explicitly (“preserve my warm tone”) or implicitly through a generic command (“improve this English”). Third, the model produces a candidate reformulation shaped by training distributions, system constraints, prompt wording, and interface conditions. Fourth, the writer may accept, reject, combine, or reprompt, thereby determining whether the generated norm enters the final human-facing text. Agency is therefore distributed but not symmetrical: the model has substantial transformational capacity, whereas responsibility for communicative uptake remains human. The empirical question is what kinds of transformations become visible when the prompt provides more or less information about audience and identity.

3. Methodology

3.1 Research design and epistemological scope

The study was situated at Taif University in Taif, Saudi Arabia, within an English-language higher-education context in which multilingual students engage with English for academic and professional purposes. Taif University provides an institutionally relevant setting for the study because English-language education and technology-mediated learning form part of its academic environment. The institutional location is treated here as the research setting rather than as a source of participant data: the present controlled elicitation phase did not recruit students or collect identifiable learner texts. Accordingly, the findings should not be generalized to Taif University students as a population.

The study adopts a controlled qualitative linguistic elicitation design. Controlled elicitation is appropriate because the research question concerns transformation rather than population prevalence: the analytic object is the relation between a source utterance and alternative AI-mediated reformulations under contrasting prompt conditions. The design is comparable to an algorithmic audit at the level of discourse, but its purpose is interpretive rather than benchmark-oriented. It examines what linguistic values become observable in outputs and how those values relate to established constructs in pragmatics, discourse analysis, and sociolinguistics.

The study makes no inference about model cognition or intention. Outputs are treated as textual events produced under specified conditions. Nor are the source utterances represented as authentic learner errors. They are researcher-designed stimuli that instantiate analytically relevant phenomena. This distinction is methodologically important: claims are restricted to textual transformation patterns within the bounded elicitation corpus.

3.2 Research setting: Taif University

The research setting was Taif University, Taif, Kingdom of Saudi Arabia. The study is positioned within the university’s English-language and digitally mediated educational environment. This contextualization is especially relevant to an applied-linguistic investigation of GenAI rewriting because the focal phenomena—English academic discourse, pragmatic appropriateness, multilingual identity, culturally situated formulaicity, and audience-sensitive revision—are salient in Saudi tertiary English-language contexts.

The institutional setting also provides a principled rationale for including Arabic-influenced English, code-meshing, culturally elaborated relational formulas, and multilingual identity among the elicitation phenomena. These features were selected to model linguistically plausible issues in Arabic-English higher-education communication; however, the 12 source utterances remain researcher-designed stimuli and must not be represented as quotations from Taif University students. No demographic claims, proficiency distributions, participant numbers, or learning outcomes are inferred from the controlled corpus.

This distinction between institutional setting and participant sampling is maintained throughout the article. Taif University anchors the sociolinguistic ecology in which the research problem is formulated, whereas the empirical object in the present phase is the transformation produced by GenAI under contrasting prompt conditions. A subsequent participant-based phase could recruit Taif University EFL students, subject to institutional ethical approval and informed consent, and triangulate authentic drafts, prompt histories, screen recordings, and stimulated-recall interviews.

3.3 Corpus construction and prompt conditions

The corpus contains 12 source utterances constructed to instantiate phenomena recurrent in applied-linguistic research: direct requests, epistemic stance, Arabic-influenced formulaicity, interpersonal solidarity, pluricentric professional lexis, code-meshing, disagreement, academic self-mention, informal evaluative stance, idiomatic transfer, relational deference, and explicit multilingual identity. Each source was paired with two revision conditions. Condition A used the underspecified instruction “Improve this English.” Condition B specified the projected audience and required preservation of relevant stance, cultural positioning, or multilingual voice. The contrast operationalizes a central distinction in AMAD: default normative mediation versus audience-conditioned mediation.

All cases were generated in September 2026 in the same ChatGPT model environment (GPT-5.6 Sol). The unit of analysis was the source-output transformation, not the source sentence in isolation. The dataset therefore comprised 24 model outputs organized into 12 matched case sets. Because the corpus was researcher-designed and model-generated, it contained no human participants, personal data, or claims about learning outcomes. Model/version and date are reported because outputs from generative systems are temporally contingent and may not be exactly reproducible after system updates.

3.4 Analytic framework and audit trail

Analysis combined deductive functional categories with inductive refinement. The initial codebook distinguished lexical normalization, grammatical standardization, mitigation/intensification, epistemic recalibration, register shift, explicitation, cultural-pragmatic substitution, code-meshing retention/deletion, self-mention redistribution, and audience-oriented reformulation. Each transformation was then located at one or more linguistic levels: lexicogrammar, pragmatics, discourse, and sociolinguistics. Coding was relational: a category was assigned only when a contrast between source and output licensed the interpretation.

A three-step audit trail strengthened interpretive transparency. First, each case was coded independently at the micro-linguistic level by identifying changed lexical, grammatical, or discourse resources. Second, the functional consequence of each change was recorded, for example reduced entitlement, weaker epistemic commitment, increased formality, or loss of a multilingual index. Third, generic and audience-aware outputs were compared within the same case to identify prompt-conditioned divergence. Higher-order themes were generated only after this within-case comparison. Deviant cases—instances in which audience-explicit prompting did not preserve markedness or where generic revision retained it—were retained rather than forced into the dominant pattern.

No inter-rater reliability statistic is reported because the study did not employ a second independent coder. Instead of presenting pseudo-precision, the analysis uses visible source-output pairs, an explicit codebook, within-case contrasts, and bounded claims as forms of analytic accountability. This remains a limitation. A participant-based extension should add independent coding, stimulated recall, prompt histories, and screen recordings to triangulate textual transformation with writers’ interpretations.

3.5 Data display

Table 1. Controlled elicitation corpus and principal analytic coding.

Phenomenon	Source utterance	Generic revision	Audience/identity-aware revision	Primary code
Direct request	Send me the report today, please.	Could you please send me the report today?	Would you please send me the report today? I would appreciate it.	Mitigation; interpersonal stance
Epistemic stance	I think this result is definitely because of motivation.	This result is likely due to motivation.	I think this result may be related to motivation, although the evidence should be interpreted cautiously.	Epistemic recalibration
Arabic-influenced formulaicity	I hope this email finds you in the best health and happiness.	I hope this email finds you well.	I hope this email finds you well. [Retains the warm relational opening while making it conventional in English.]	Formulaic compression
Solidarity	My dear professor, thank you from my heart for your support.	Dear Professor, thank you very much for your support.	Dear Professor, thank you sincerely for your generous support; I deeply appreciate it.	Affect attenuation / retention
Global Englishes	Kindly revert to me when you receive the document.	Please reply when you receive the document.	Kindly let me know when you receive the document.	Pluricentric lexical normalization
Code-meshing	The meeting was, wallahi, more difficult than I expected.	The meeting was honestly more difficult than I expected.	The meeting was, wallahi (truly), more difficult than I expected.	Deletion vs glossed retention
Direct disagreement	I do not agree with this point because it ignores the context.	I disagree with this point because it overlooks the context.	I respectfully disagree with this point because it appears to overlook the relevant context.	Politeness and hedging
Academic self-mention	I strongly believe my analysis proves that identity is changing.	The analysis demonstrates that identity is changing.	I argue that the analysis provides evidence that identity is changing.	Authorial voice redistribution
Informal digital stance	This finding is super important and honestly kind of shocking.	This finding is highly significant and somewhat surprising.	This finding is particularly important and, from the writer's perspective, unexpectedly strong.	Register elevation
Idiomatic transfer	He put the points on the letters.	He clarified the issue.	He clarified the issue explicitly, preserving the source expression's emphasis on making every detail clear.	Cultural-pragmatic substitution

Relational request	If you do not mind, doctor, I need your help in this matter.	Could you please help me with this matter?	If you do not mind, Professor, I would greatly appreciate your help with this matter.	Role address; deference
Multilingual identity	As an Arab student, I want my English to sound like me, not like a machine.	As an Arab student, I want my English to reflect my own voice rather than sound machine-generated.	As an Arab student, I want my English to retain my own voice and cultural positioning rather than being standardized into machine-like prose.	Identity preservation

Table 1. Controlled elicitation corpus and principal analytic coding.

4. Results and Discussion

4.1 Generic improvement as sociolinguistic normalization

The strongest cross-case pattern was movement toward an implicitly unmarked professional-English norm. This process was visible even when the source was intelligible and pragmatically functional. “Kindly revert to me when you receive the document,” for example, became “Please reply when you receive the document.” The revision is readily defensible for some audiences, but it also reclassifies a pluricentric professional expression as less suitable than a form associated with wider Anglo-American institutional circulation. Similarly, the Arabic-influenced relational opening “I hope this email finds you in the best health and happiness” was compressed to “I hope this email finds you well.” The change reduces formulaic elaboration and increases conformity to a routinized Anglophone email opening.

These examples show why error-correction terminology is insufficient. Neither transformation is primarily grammatical. What changes is markedness. The generic prompt supplies no audience, variety, or identity information, so the system effectively infers a default norm. AMAD describes this as default normative mediation: absent explicit constraints, the model resolves uncertainty by producing a form with high institutional recognizability. The sociolinguistic consequence is a subtle hierarchy in which dominant conventions appear neutral while culturally or regionally situated alternatives appear in need of correction.

A deviant pattern complicates any claim that the model invariably erases variation. Some marked meanings were semantically retained through paraphrase. The idiomatic-transfer case, “He put the points on the letters,” became “He clarified the issue.” The communicative proposition survives, but the figurative cultural trace disappears. This distinction between propositional preservation and indexical preservation recurred across the corpus and is analytically important: an output may be semantically adequate while sociolinguistically reductive.

4.2 Pragmatic mitigation and the redistribution of interpersonal force

Requests and disagreements were repeatedly reformulated through conventional mitigation. “Send me the report today, please” became “Could you please send me the report today?”, while “I do not agree with this point because it ignores the context” became “I disagree with this point because it overlooks the context.” In the audience-aware condition, disagreement was further recast as “I respectfully disagree” and “appears to overlook.” These changes preserve the broad speech act while modifying entitlement, face management, and the projected relationship between interlocutors.

The pattern demonstrates that GenAI acts as a pragmatic mediator. Modal interrogatives, respectful adverbials, appreciation formulas, and evidential softeners are not merely more formal wordings; they recalibrate the illocutionary force of the utterance. This matters pedagogically because a learner may intend urgency, solidarity, or principled disagreement and yet receive a revision that changes the interpersonal stance. Treating such output as automatically “more polite” risks collapsing pragmatic appropriateness into maximal mitigation.

Audience-explicit prompting produced greater sensitivity to relational meaning. In the deference case, the generic output removed the role address and conditional preface, whereas the audience-aware output retained “If you do not mind, Professor” and added appreciation. The contrast supports RQ3: specifying audience relationship can constrain generic normalization, although it may also intensify politeness beyond the source. Audience information therefore does not simply preserve meaning; it becomes an additional resource from which the model constructs a social relation.

4.3 Epistemic recalibration and authorial visibility

A third process involved changes to commitment and authorial presence. The source “I think this result is definitely because of motivation” combined explicit self-mention with strong causal certainty. The generic revision, “This result is likely due to motivation,” removed the first-person frame and reduced certainty. The audience-aware revision, “I think this result may be related to motivation, although the evidence should be interpreted cautiously,” restored self-positioning but weakened causal force further. Each version occupies a different epistemic position.

The academic self-mention case revealed a related redistribution. “I strongly believe my analysis proves that identity is changing” became “The analysis demonstrates that identity is changing” under generic improvement. The revision removes the writer while retaining a strong evidential verb. By contrast, the audience-aware version—“I argue that the analysis provides evidence that identity is changing”—maintains authorial responsibility while moderating the knowledge claim. The latter is not simply a more elegant sentence; it reorganizes the relation among writer, evidence, and proposition.

These transformations answer RQ2 by showing that AI revision can alter discourse-level responsibility. In academic writing, voice is partly realized through precisely these choices: who claims, how strongly, on what evidential basis, and with what degree of interpersonal visibility. Critical AI literacy must therefore include an epistemic audit of generated revisions. Writers need to ask not only whether a sentence is fluent but whether the revised commitment accurately represents what the evidence licenses and what the author is prepared to claim.

4.4 Multilingual indexicality and prompt-sensitive preservation

The code-meshing and multilingual-identity cases most clearly exposed the difference between semantic and indexical meaning. In the source “The meeting was, wallahi, more difficult than I expected,” the generic revision replaced wallahi with “honestly.” The proposition remains intelligible, but a culturally recognizable resource is converted into a monolingual English stance marker. Under audience-aware instruction, wallahi was retained and glossed as “truly.” This solution increases accessibility without requiring deletion.

The contrast should not be romanticized as evidence that all multilingual forms must always be preserved. Communicative appropriateness depends on audience knowledge, genre, institutional expectations, and writer purpose. The important finding is that generic improvement made the decision implicitly, whereas audience-explicit prompting opened a larger design space. The writer could preserve, gloss, translate, or replace the form strategically. In AMAD terms, explicit audience specification redistributed agency back toward the writer by making the desired indexical effect part of the prompt.

The explicit identity statement produced a similar result. Both outputs preserved the speaker’s identification as an Arab student, but the audience-aware version foregrounded “voice and cultural positioning” and explicitly resisted standardization into “machine-like prose.” This case is partly metalinguistic, yet it demonstrates the model’s responsiveness to identity constraints. The finding is compatible with translanguaging perspectives that treat multilingual resources as part of a speaker’s integrated repertoire rather than noise to be removed.

4.5 Register elevation and affective flattening

Generic revision also tended to elevate register by substituting conventional academic or professional lexis for conversational evaluation. “Super important” became “highly significant,” and “honestly kind of shocking” became “somewhat surprising.” The transformation increases formal congruence but weakens immediacy and affect. Likewise, “thank you from my heart” became “thank you very much,” reducing culturally inflected warmth. The audience-aware version recovered some affect through “thank you sincerely” and “I deeply appreciate it.”

Register shift therefore cannot be treated as a one-dimensional movement from informal to appropriate. Register choices index persona and relationship. A more institutionally conventional form may be desirable in a research article and undesirable in a message whose communicative purpose depends on warmth or solidarity. Across the corpus, the model’s generic revisions repeatedly privileged institutional fit over expressive markedness. Audience-conditioned prompting partially moderated this tendency, suggesting that effective human-AI writing depends on specifying the social work the text must perform.

4.7 From audience design to algorithmically mediated audience design

The findings extend audience design in a specific way. Bell’s framework explains style-shifting as orientation to socially differentiated audiences. In AI-assisted writing, however, the path from writer to audience may contain a generative intermediary that transforms the linguistic resources through which that orientation is realized. The writer designs for the eventual addressee,

but must also design a prompt that instructs the model how to preserve or recalibrate that relationship. Audience design consequently becomes recursive.

This recursion does not require anthropomorphizing the model. The system is not a ratified human participant with intentions or social membership. Its relevance is interactional and infrastructural: writers orient to anticipated model behavior, and model outputs alter the text that human audiences receive. This formulation aligns with the current Applied Linguistics agenda of using AI to revisit foundational questions of meaning-making and interaction rather than treating model performance as an end in itself. It also complements Tang's (2026) AI-textuality by locating one consequential site of human-AI textuality in the mediation of audience-oriented style.

4.8 "Improvement" as covert norm selection

The empirical contrast also exposes the normative content of generic improvement. When a prompt supplies no explicit target variety, audience, or identity constraint, the model must nevertheless produce a candidate norm. Across these cases, that norm tended toward broadly recognizable professional or academic English, with reduced regional markedness, increased mitigation, and elevated register. Such revisions can be useful. The problem arises when their normativity is rendered invisible by the label "improvement."

Applied linguistics has long shown that correctness, appropriateness, and legitimacy are socially organized rather than purely formal properties. GenAI operationalizes these judgments at scale. A recurrent suggestion can become a form of de facto language policy even without an explicit institutional rule. This is why concerns about digital linguistic imperialism are not external ethical add-ons; they are directly linguistic. The central question is whose pragmatic and stylistic conventions are encoded as defaults and what happens to users whose repertoires diverge from them.

4.9 Voice, agency, and human oversight

The findings further suggest that agency is best understood as the capacity to regulate transformation rather than simply the choice to use or not use AI. Audience-aware prompts preserved more relational and multilingual meaning because they imposed constraints on the rewriting task. This supports recent work showing that advanced multilingual writers vary in whether they treat GenAI as a prescriptive tool, dialogic partner, or distributed component of writing activity (Wang et al., 2026). It also converges with Khuder's (2026) emphasis on feedback-seeking as a means of protecting disciplinary voice.

Human oversight, on this account, is linguistically substantive. It requires writers to recognize shifts in stance, entitlement, evidentiality, register, self-mention, and indexicality. A writer who cannot identify those dimensions may formally retain decision-making authority while functionally delegating consequential choices to the system. Pedagogical approaches should therefore move beyond prompt engineering as an instrumental skill and teach comparative linguistic diagnosis: source, AI revision, functional change, audience consequence, and final authorial decision.

4.10 Implications for multilingual and L2 writing pedagogy

For L2 and multilingual writing, the analysis supports a contrastive pedagogy of AI revision. Learners can compare an original utterance with multiple generated alternatives and annotate what changes at four levels: lexicogrammar, pragmatics, discourse, and sociolinguistics. Instead of asking "Which version is correct?", instruction can ask "Which version best enacts the intended relationship and why?" Such tasks connect language awareness to communicative purpose and make AI output an object of analysis rather than an answer key.

A second implication concerns translanguaging and Global Englishes. Teachers can ask learners to test whether AI preserves regionally meaningful expressions under different audience descriptions, then discuss when glossing, accommodation, or retention is appropriate. This reframes standardization as one rhetorical option among several. The objective is not resistance to all editing, but informed control over the social meanings that editing can redistribute.

Third, academic writing instruction should foreground epistemic consequences. Learners should track changes to hedges, boosters, causal verbs, self-mention, evidentials, and reporting verbs. An AI-generated sentence that sounds more academic may overstate evidence, erase responsibility, or impose a disciplinary stance the writer does not endorse. This is particularly important for novice scholars who may equate lexical sophistication with epistemic appropriateness.

4.6 Methodological Contribution, Limitations, and Research Agenda

Methodologically, the study demonstrates a transparent way to use controlled GenAI outputs as applied-linguistic data without confusing them with participant data. The design specifies the model environment and generation period, defines the transformation pair as the unit of analysis, contrasts prompt conditions, displays the complete small corpus, and limits inference

to textual behavior. These choices respond to the reproducibility problem created by model updates: exact replication may be impossible, but analytic traceability can still be strengthened through explicit documentation.

The design also has clear limitations. Twelve researcher-designed cases cannot establish the prevalence of a pattern across models, languages, genres, or user populations. The stimuli inevitably reflect the researcher's theoretical interests. Only one model environment and one temporal snapshot were examined, and no independent coder was available. Most importantly, because there were no human participants, the study cannot show how actual multilingual writers interpret, accept, reject, or learn from these revisions. The findings should therefore be read as an empirically illustrated theoretical intervention rather than a population-level evaluation of GenAI.

A stronger next phase at Taif University would combine algorithmic auditing with participant-based qualitative inquiry, subject to the university's applicable research-ethics procedures. Multilingual EFL students could complete authentic academic and professional writing tasks while screen recordings capture prompts, revisions, and acceptance decisions. Stimulated-recall interviews could then examine why writers accepted or rejected particular changes. Independent coding of textual transformations could be combined with participant accounts of voice and audience. Cross-model and cross-language comparisons would establish whether observed normalization is platform-specific or infrastructurally recurrent.

The AMAD framework generates testable propositions for that research. Generic prompts should produce greater convergence toward dominant professional norms than audience-explicit prompts; identity-explicit constraints should increase retention or glossing of multilingual resources; writers with stronger metapragmatic awareness should reject a greater proportion of revisions that alter intended stance; and institutional genres with highly codified norms may produce different patterns from relational or public-facing genres. These propositions move the framework beyond description and provide a basis for comparative empirical work.

5. Conclusion

Generative AI has made linguistic mediation newly visible. A user can now watch a system transform directness into mitigation, certainty into caution, self-mention into impersonality, multilingual markedness into standardized English, or informal affect into institutional register within seconds. These transformations are often useful, but they are never linguistically empty. They reorganize the resources through which writers enact relationships, claim knowledge, project identities, and orient to audiences.

This study has conceptualized that process as algorithmically mediated audience design. The controlled corpus showed five recurrent tendencies—sociolinguistic normalization, pragmatic mitigation, epistemic recalibration, redistribution of authorial voice, and selective preservation of multilingual indexicality—and demonstrated that explicit audience and identity instructions can alter, though not abolish, these tendencies. The theoretical implication is that audience design in AI-mediated communication is recursive: writers orient to human recipients through an algorithmic intermediary whose transformations become part of the communicative event.

For applied linguistics, the central challenge is therefore not to decide whether GenAI is intrinsically beneficial or harmful. It is to develop concepts and methods capable of identifying what linguistic work these systems perform, whose norms they privilege, and how human users can retain informed control over consequential choices. An applied-linguistic approach to AI should evaluate not only fluency and accuracy, but also voice, pragmatic force, epistemic responsibility, multilingual legitimacy, and the distribution of communicative agency.

Funding: This research received no external funding.

Conflicts of Interest: The authors declare no conflict of interest.

ORCID iD: First Author ORCID, <https://orcid.org/0009-0004-2367-5737> ; Second Author ORCID, <https://orcid.org/0000-0001-8673-1100>

Ethics Statement: The controlled elicitation phase used researcher-designed stimuli and model-generated outputs and did not recruit human participants or collect identifiable learner data.

Data Availability Statement: The complete controlled elicitation corpus is reported in Table 1. Additional prompt/output records may be supplied as supplementary material.

Generative AI Disclosure: Generative AI outputs constitute the object of analysis in this study. The authors remain responsible for the analysis, interpretation, and final manuscript.

References

- [1]. Alvero, A. J., Sedlacek, Q., León, M., & Peña, C. (2025). Digital accents, homogeneity-by-design, and the evolving social science of written language. *Annual Review of Applied Linguistics*, 45, 50–68. <https://doi.org/10.1017/S0267190525000042>
- [2]. Bakhtin, M. M. (1981). *The dialogic imagination: Four essays* (M. Holquist, Ed.; C. Emerson & M. Holquist, Trans.). University of Texas Press.
- [3]. Bell, A. (1984). Language style as audience design. *Language in Society*, 13(2), 145–204.
- [4]. Bell, A. (1991). *The language of news media*. Blackwell.
- [5]. Bell, A. (2001). Back in style: Reworking audience design. In P. Eckert & J. R. Rickford (Eds.), *Style and sociolinguistic variation* (pp. 139–169). Cambridge University Press.
- [6]. Darvin, R. (2025). Identity and investment in the age of generative AI. *Annual Review of Applied Linguistics*, 45, 10–27. <https://doi.org/10.1017/S0267190525100135>
- [7]. Khuder, B. (2026). Enhancing disciplinary voice through feedback-seeking in AI-assisted doctoral writing for publication. *Applied Linguistics*, 47(2), 372–387. <https://doi.org/10.1093/applin/amaf022>
- [8]. Lee, S. (2025). Generative AI and its dilemmas: Exploring AI from a translanguaging perspective. *Applied Linguistics*, 46(4), 709–717. <https://doi.org/10.1093/applin/amaf049>
- [9]. Martinez, R., Martinez, G., & Curle, S. (2026). Shaping an enduring conversation on AI in applied linguistics. *Applied Linguistics*, 47(2), 191–192. <https://doi.org/10.1093/applin/amag050>
- [10]. Mohammad Hosseinpour, R. (2026). Digital linguistic imperialism: Can we decolonize English when the AI writing tools are colonial? *Applied Linguistics*, advance article, amag073. <https://doi.org/10.1093/applin/amag073>
- [11]. O’Keeffe, A. (2006). *Investigating media discourse*. Routledge.
- [12]. Tang, K.-S. (2026). AI-textuality: Expanding intertextuality to theorize human-AI interaction with generative artificial intelligence. *Applied Linguistics*, 47(2), 353–371. <https://doi.org/10.1093/applin/amaf016>
- [13]. UNESCO. (2023). *Guidance for generative AI in education and research*. UNESCO.
- [14]. Wang, C., et al. (2026). Advanced multi-lingual writers’ self-directed use of generative AI in academic writing: Rethinking writing, authorship, and learning. *Applied Linguistics*, 47(2), 388–405. <https://doi.org/10.1093/applin/amaf057>