

VOL. XI ISSUE I, WINTER (March-2026)

GER

GLOBAL ECONOMICS REVIEW

GLOBAL ECONOMICS REVIEW (GER)

Title: Fluent Access or Reliable Access? Artificial Intelligence, Multilingual Customer Communication, and the Politics of Understanding

Abstract

While AI enables public and private organizations to translate, simplify, and tailor messages for multilingual audiences at scale, the social, cognitive, and governance implications of AI-assisted language transformation remain under-theorized. Grounded in AI-mediated communication (AI-MC), organizational language, trust theory, and responsible governance frameworks, this article argues that superficial linguistic fluency alone does not confer institutional value. The research distinguishes between Fluent Access (smooth, natural, and readable text) and Reliable Access (communication that preserves meaning, supports proper action, conveys respect, and remains rooted in human oversight). We demonstrate that unmediated AI translation creates an "Overconfidence Gap" (+32.3 percentage points), where users express high confidence despite misunderstanding unclear obligations. Crucially, integrating plain-language source texts, competent bilingual review, and accessible human intervention pathways closes this confidence gap and establishes enduring institutional trust.

Keywords: AI-Mediated Communication (AI-MC), Multilingual AI Translation, Overconfidence Gap, Reliable Access, Responsible AI Governance

Authors:

Sana Hussain: (Corresponding Author)

PhD Scholar, Business Administration, The University of the Potomac, Virginia, United States (US)
(Email: sanahusan143@gmail.com)

Pages: 1-13

DOI: 10.31703/ger.2026(XI-I).01

DOI link: [https://dx.doi.org/10.31703/ger.2026\(XI-I\).01](https://dx.doi.org/10.31703/ger.2026(XI-I).01)

Article link: <https://gerjournal.com/article/fluent-access-or-reliable-access-artificial-intelligence-multilingual-customer-communication-and-the-politics-of-understanding>

Full-text Link: <https://gerjournal.com/article/fluent-access-or-reliable-access-artificial-intelligence-multilingual-customer-communication-and-the-politics-of-understanding>

Pdf link: <https://www.gssrjournal.com/jadmin/Author/31rvIolA2.pdf>

Global Economics Review

p-ISSN: 2521-2974 e-ISSN: 2707-0093

DOI(journal): 10.31703/ger

Volume: XI (2026)

DOI (volume): 10.31703/ger.2026(XI)

Issue: I Winter (March-2026)

DOI(Issue): 10.31703/ger.2026(XI-I)

Home Page

www.gerjournal.com

Volume: XI (2026)

<https://www.gerjournal.com/Current-issue>

Issue: I-Winter (March 2026)

<https://www.gerjournal.com/issue/11/1/2026>

Scope

<https://www.gerjournal.com/about-us/scope>

Submission

<https://humaglobe.com/index.php/ger/submissions>

Scan the QR to visit us



Google
scholar



Citing this Article

Article Serial	01
Article Title	Fluent Access or Reliable Access? Artificial Intelligence, Multilingual Customer Communication, and the Politics of Understanding
Authors	Sana Hussan
DOI	10.31703/ger.2026(XI-I).01
Pages	1-13
Year	2026
Volume	XI
Issue	I
Referencing & Citing Styles	
APA	Hussan, S. (2026). Fluent Access or Reliable Access? Artificial Intelligence, Multilingual Customer Communication, and the Politics of Understanding. <i>Global Economics Review</i> , 11(1), 1-13. https://doi.org/10.31703/ger.2026(XI-I).01
CHICAGO	Hussan, Sana. 2026. "Fluent Access or Reliable Access? Artificial Intelligence, Multilingual Customer Communication, and the Politics of Understanding." <i>Global Economics Review</i> 11 (1):1-13. doi: 10.31703/ger.2026(XI-I).01.
HARVARD	HUSSAN, S. 2026. Fluent Access or Reliable Access? Artificial Intelligence, Multilingual Customer Communication, and the Politics of Understanding. <i>Global Economics Review</i> , 11, 1-13.
MHRA	Hussan, Sana. 2026. 'Fluent Access or Reliable Access? Artificial Intelligence, Multilingual Customer Communication, and the Politics of Understanding', <i>Global Economics Review</i> , 11: 1-13.
MLA	Hussan, Sana. "Fluent Access or Reliable Access? Artificial Intelligence, Multilingual Customer Communication, and the Politics of Understanding." <i>Global Economics Review</i> 11.1 (2026): 1-13. Print.
OXFORD	Hussan, Sana (2026), 'Fluent Access or Reliable Access? Artificial Intelligence, Multilingual Customer Communication, and the Politics of Understanding', <i>Global Economics Review</i> , 11 (1), 1-13.
TURABIAN	Hussan, Sana. "Fluent Access or Reliable Access? Artificial Intelligence, Multilingual Customer Communication, and the Politics of Understanding." <i>Global Economics Review</i> 11, no. 1 (2026): 1-13. https://dx.doi.org/10.31703/ger.2026(XI-I).01 .

Fluent Access or Reliable Access? Artificial Intelligence, Multilingual Customer Communication, and the Politics of Understanding



Sana Hussan (Corresponding Author)¹

¹ PhD Scholar, Business Administration, The University of the Potomac, Virginia, United States (US)
(Email: sanahusan143@gmail.com)

Abstract

While AI enables public and private organizations to translate, simplify, and tailor messages for multilingual audiences at scale, the social, cognitive, and governance implications of AI-assisted language transformation remain under-theorized. Grounded in AI-mediated communication (AI-MC), organizational language, trust theory, and responsible governance frameworks, this article argues that superficial linguistic fluency alone does not confer institutional value. The research distinguish between Fluent Access (smooth, natural, and readable text) and Reliable Access (communication that preserves meaning, supports proper action, conveys respect, and remains rooted in human oversight). We demonstrate that unmediated AI translation creates an "Overconfidence Gap" (+32.3 percentage points), where users express high confidence despite misunderstanding unclear obligations. Crucially, integrating plain-language source texts, competent bilingual review, and accessible human intervention pathways closes this confidence gap and establishes enduring institutional trust.

Keywords: *AI-Mediated Communication (AI-MC), Multilingual AI Translation, Overconfidence Gap, Reliable Access, Responsible AI Governance*

Introduction

The prospect of using AI to make communication easier is an exciting institutional remedy to a longstanding and entrenched problem in the governance of the modern state: the capacity to distill intricate information and phrases into simple ones, using language that is accessible to diverse language groups, at speeds and costs we could never have imagined. This promise is especially enticing in multilingual societies and in economies that rely on digital service delivery. There are growing regulatory requirements and ethical considerations for public health agencies, electric and water utilities, higher education financial aid offices, retail banking service providers, and social safety net programs to reach patients as well as constituents whose home languages vary and are often different from English, as well as those whose literacy rates are under 50 percent and whose understanding of complex safety administration jargon is limited (Hancock, Naaman, & Levy, 2020; Wang, 2024). When it comes to limitless scalability, where traditional human translation and bilingual staff have historically been operationally limited, slow, and expensive, generative artificial intelligence systems such as neural machine translation (NMT) and large language model (LLM) agents seem to be the answer.

At the same time, however, an urgent question is emerging about how such technology can be integrated uncritically into institutional communication technologies, serving as a point of departure for business communicators, public administrators, and AI governance scholars. Does AI construe real communicative access for multilingual users, or can it often only be seen as superficial access, in the form of prose which superficially appears fluent, but substantively is ungrounded and unreliable? An automated computational



agent can translate a complicated regulatory notice, medical treatment regimen, or financial repayment term into a target language, changing much more than simply surface words. It must actively adjust the text's inherent pragmatic cues, legal force, condition qualifications, and administrative requirements.

This research article argues that readability tools, the machine translation BLEU score, or an initial user satisfaction rating are not the only metrics for evaluating multilingual artificial intelligence communication. It can be flawlessly written, easy for a reader with limited English to read, and perfect for the recipient, and yet fail to accurately report a legally binding deadline, minimize or omit a warning in a critically important law, miss out on significant honorifics or important cultural cues, or allow a reader to unconsciously assume they know how to read it in such a way as to give them false assurance that they are correct. In formal institutional communications, this failure can translate into immediate, severe real-world consequences. Whether it is medical interventions, housing assistance renewals, utilities shutoffs, coming up to date on debt, or making legal appeals to cite, recipients of messages depend on institutional messages to make critical life decisions. Fluency, even when untrustworthy, might increase distrust of administrative words or messages, leading vulnerable individuals to take rash action; a computer system that creates fluent but unreliable text might exacerbate social and economic divisions.

The fundamental issue is not only linguistic and technical; it is also organizational, behavioral, and political. Language at one level of formal institutions is always an interpretive means of transmission. Instead, language is inextricably linked to institutional authority, social standing, citizenship engagement, and institutional legitimacy (Neeley, 2013; Tenzer, Pudelko, & Harzing, 2014). Public or business communications that require citizens only to speak in dense, technical, or dominant-language forms impose a significant cognitive, financial, and emotional 'administrative burden' on those who don't speak that language as their first language. Generative AI tools can alleviate this administrative pain, but ungoverned AI deployment often masks underlying ambiguity behind an attractive form with highly persuasive content.

This article proposes a full-picture, communication-centric framework to overcome this significant theoretical and practical concern, based on the theoretical differences between Fluent Access and Reliable Access. Based on the macro-empirical analyses of N = 12,450 observations by consumer evaluators in 5 key institutional service categories, we examine the structural causal pathways, psychological mechanisms, and organizational structures that distinguish shallow, smooth communication from deep and trustworthy communication.

Literature Review:

Language Asymmetries, Status Loss, and Institutional Power

There is subordination of authority and stratification of society in language within formal organizations such as institutions and governmental entities. Held up as a benchmark in organizational studies (Neeley, 2013), institutional language mandates create stark power disparities and heighten status anxieties, competence loss, and psychological defensiveness among those who do not speak the dominant group's language. When people are exposed to institutions that operate in a language other than their native one, they face increased cognitive burden and emotional vulnerability in the face of communication failures, which are likely to be attributed to their own lack of ability and read as the service's or institution's inability to communicate.

Extending this idea to communication with customers or the public, Tenzer, Pudelko, and Harzing (2014) found that language barriers systematically affect the ability to develop trust. Clearer messages signal more integrity and benevolence in business or public service. Dense, legalese, and/or culturally insensitive communication puts the non-native person or entity in the institution's communication in a position to infer a lack of institutional respect. Multilingual recipients face a double administrative burden: they must look up unfamiliar terms while also navigating unspoken organizational norms, potential conditional clauses, and the risk of statutory penalties.

Generative AI seems to be leveling the playing field for accessing institutions; instant translation and text-simplification features are highly beneficial. But as Wang (2024) points out, ungoverned machine

translation often replaces explicit language barriers with subtle, opaque semantic errors. If user-institutional interaction is cut short or glossed over in translating institutional text to ease the interaction burden, little to nothing will be achieved.

AI-Mediated Communication (AI-MC) and Algorithmic Agency

More recently, Hancock, Naaman, and Levy (2020) defined AI-Mediated Communication (AI-MC) as communication supported by algorithms that leverage computation to modify, augment, or create a specific natural language input on behalf of a communicator for a specific relationship or operational purpose. AI agents are more than passive communication channels in multilingual customer communication. They are active, communicative mediators that change message framing, syntactic complexity, emotional flavor, and other attributes.

AI-MC is, in multilingual institutional contexts, a threefold process: (1) a syntactic operation of translation from one language system to another; (2) a semantic abstraction or text simplification work; and (3) an operation of aligning the text to the norms of the target group. There are different failure modes per operation. Monitoring the accuracy of machine translation can reveal a tendency by a machine to change a 'must' in the original document to a 'should' in the translated version; text simplification will often strip out conditional sub-clauses; and automatic summarization will often take the 'must' from the lower-level ground rules and elevate them to the top-level 'claim'. Non-native recipients are often not fluent enough to verify the target text against the source documents, and in this situation they lack any baseline verification capacity, resulting in acute cognitive over-confidence (Schilke & Reimann, 2025).

The Dual-Pillar Framework: Fluent Access vs. Reliable Access

Since there are conceptual ambiguities about “automated communication,” we introduce the dual-pillar approach of Fluent Access / Reliable Access. Fluent Access is the surface-level communicative performance; what we see here is reading ease, syntactic naturalness, grammatical correctness, and minimal immediate processing friction. Reliable Access means that the system-level communicative integrity (i.e., the exactness of semantic preservation, actionable correctness of statements, context-appropriate pragmatic tone, and inclusion in an accessible organization of communicative recourse) needs to be addressed.

What's shown to be sensitive is the separation of these two pillars in today's use of AI. While state-of-the-art LLMs can provide near-perfectly fluent access, an organization can still fail on the same metrics to deliver reliable access. A utility chatbot using automation can produce a perfect, caring message in the target language but then give an inaccurate answer to a low-income customer about whether a shut-off notice has been delayed when, in fact, it has not been, and the GOC must be filed. In cases of limited access where no real access can be achieved, communication is used to fulfill the role of accessibility and helps to raise the risk for the customer.

Moreover, this separation reshapes how consumer satisfaction indicators are understood and applied in AI governance. A first impression of high satisfaction scores when customer service information is initially translated into English via AI may reflect the translated text's readability and fast posting times, but not necessarily accurate comprehension of the content. Focusing only on “here and now” sentiment or CSAT scores can lead organizations to accept systems that may consistently misrepresent the information their least powerful constituents need.

Administrative Burden, Cognitive Friction, and Algorithmic Vulnerability

According to the literature on administrative burden, there are three types of administrative burden associated with contact with the public and/or with customers, namely learning costs (time it takes to learn about the rules governing service, as well as who qualifies and who does not), compliance costs (time it takes to submit the necessary paperwork and documentation), and psychological costs (feeling anxious, stigmatized, and loss of control). Administrative burden is magnified when interacting with multilingual

services. For non-native speakers, this extra cognitive energy goes into understanding institutional vocabulary.

AI translation can reduce learning costs and deliver information in the target language, but the machine's output can often increase compliance and psychological costs. Finally, if an AI translation produces an ambiguous sentence or omits a necessary instruction, serious compliance consequences can follow such as late applications or failure to receive rewards, like not being selected for a grant—often without anyone understanding that the translation was the issue. This creates what might be termed as 'algorithmic vulnerability' where users with fewer language resources are placed at a greater risk by automated errors.

Algorithmic vulnerability also affects access to administrative remedies, as this is unequal. People who speak the dominant language can easily reach out to customer service if they spot an error in the automated message, explain the problem, and ask for it to be addressed. But for non-native English speakers, challenging an incorrect AI-generated notice letter is daunting. Without explicit target-language error-correction channels, customers shoulder the administrative labor of error correction.

Pragmatic Competence, Culturally Sensitive Phrasing, and Dignity

Communicative effectiveness has much more to do with pragmatic competence handling, adapting, and understanding social context, politeness registers, honorifics, and cultural expectations than with semantic accuracy. Tone is a key signal of regard for the institution in institutional communication. Translations that express meaning correctly, but may not be delivered in the correct manner for the recipient culture would be considered incorrect, for example, when too direct language is employed in a culture which values indirect politeness, or conversely, condescending child-like language in a text simplification process.

A high level of institutional trust is lost when recipients are treated as less than human, or as a case for a machine that disseminates information. Clearly, as Castelo et al. (2019) show, if algorithms are not socially and emotionally aware, consumers will immediately resist them. Talking to customers respectfully and using pragmatic textual design (the right honorific, empathetic phrasing, culturally aligned framing, etc.) is a key psychological precondition for building legitimate customer trust in AI-controlled systems.

Trust Theory, Institutional Legitimacy, and Calibrated Confidence

Communication trust in institutions takes two forms: trust as a rational, calculus-based view of competence and accuracy, and trust as an identification-based view of benevolence, values, shared values, and respect. Automated systems can provide high surface fluency, creating positive, calculus-based trust cues that feel highly professional. But when that articulation is not necessarily factually accurate, a substantial gap emerges between the sense of reliability in a person's presence and the systemic reliability behind it.

This disjunction undermines institutional legitimacy. Micro-institutional theory states that to gain legitimacy in an organization, it must align with societal norms and stakeholder expectations (Schilke & Reimann, 2025). Such automated responses, prancing around an institution's obligations, make the moral and cognitive credibility tumble when it is discovered they don't have to. Thus, 'calibrated confidence', the desired level of confidence in oneself, must align with the reality of the communication made by the sender, for the long-term legitimacy of the organization.

Responsible Governance, Plain Source Composition, and Human Recourse

This ongoing process of validation, risk management, and human oversight is a key part of responsible AI governance frameworks, such as the National Institute of Standards and Technology (NIST) Generative AI Profile (NIST AI 600-1, 2024). Effective governance starts upstream at source text composition in multilingual customer communications. Ambiguous or over-structured source text makes machine translation algorithms more likely to compound errors, since the input is dense. Algorithmic translation fidelity requires plain-language composition in the source language.

Downstream governance requires Human-in-the-Loop (HITL) oversight and competent bilingual reviewers, and human recourse pathways must be accessible. Automated systems lack situational judgment and cultural pragmatics, as Grigsby et al. (2025) and EAS Publisher (2026) noted. Human oversight serves as a critical check and balance. Ample recourse options (dialing out to bilingual human agents directly via phone or chat, for instance) allow recipients to contest, clarify, and correct communication issues, supporting institutional legitimacy and customer trust.

Methodology

Systematic Empirical Meta-Synthesis and Sample Architecture

To achieve the above, the following steps will be performed: 3.1 Systematic Empirical Meta-Synthesis and Sample Architecture. The steps involved in this will be:

We conducted a large-scale empirical meta-synthesis of experimental data from controlled on-road user evaluation trials, linguistic accuracy audits, and organizational governance reviews to test our conceptual framework and empirically measure the gap between Fluent Access and Reliable Access. The combined empirical dataset comprises N = 12,450 observations of consumer evaluators collected across 5 service sectors in North America and Western Europe: Healthcare Administration, Municipal Utilities, Higher Education Financial Aid, Digital Banking Services, and Social Safety Net Administration.

Table 1

Systematic Empirical Meta-Synthesis Corpus & Dataset Architecture

Service Sector	Sample Size (N)	Primary Target Languages	AI Engine / Deployment	Governance Design
Healthcare Admin.	N = 2,840	Spanish, Vietnamese	LLM Simplification & NMT	AI-Only vs. Nurse Review
Municipal Utilities	N = 2,210	Spanish, Tagalog	Automated Service Bot	Unframed vs. Recourse Link
Higher Education	N = 2,450	Chinese, Arabic	Financial Aid Summarizer	Complex vs. Plain Source
Digital Finance	N = 2,650	Spanish, Portuguese	Dispute Resolution Bot	Mandatory AI Disclosure
Social Benefits	N = 2,300	Spanish, Haitian Creole	Eligibility Assistant	Full HITL Governance Model

Experimental Tasks and Sectoral Risk Profiles

Across each of the 5 institutional sectors, standardized experimental communication stimuli were constructed representing high-stakes customer touchpoints. Evaluators were presented with translated or simplified communications under four randomized experimental governance conditions: (1) Raw AI Output (unreviewed machine translation), (2) AI + Plain Source Composition (AI translation executed on simplified source text), (3) AI + Bilingual Human Review (AI translation reviewed and edited by certified bilingual staff), and (4) Full Governance Model (Plain source + Bilingual review + Embedded human recourse link).

Table 4*Institutional Sector Task Specifications & Failure Risk Profiles*

Service Sector	Experimental Communication Stimulus	Primary Operational Failure Risk	Consequence Level
Healthcare Admin.	Post-surgical discharge instructions & dosage rules	Misunderstood medication timing or warning signs	High (Clinical Risk)
Municipal Utilities	Notice of service disconnection & hardship filing	Missed shutoff suspension deadline; utility loss	High (Essential Service)
Higher Education	Financial aid verification & appeal submission form	Missed grant deadline; forfeiture of tuition aid	Medium-High (Financial)
Digital Finance	Unauthorized transaction dispute & liability notice	Missed 10-day fraud reporting window; loss of funds	High (Financial)
Social Benefits	Annual SNAP benefit renewal & income reporting notice	Failure to submit mandatory proof; benefit termination	Critical (Safety Net)

Operationalization of Constructs and Measurement Scales

The empirical meta-synthesis operationalized 6 primary antecedent variables, 4 mediating psychological perceptions, and 2 outcome constructs using standardized, validated scales across all sector sub-samples:

1. Surface Readability / Fluency (Antecedent/Mediator): Measured via standardized Flesch-Kincaid Grade Level scores, Automated Readability Index (ARI), and a 7-point subjective user readability scale (1 = Extremely Difficult, 7 = Exceptionally Clear and Smooth).
2. Objective Comprehension Accuracy (Mediator/Outcome): Operationalized via 5-item multiple-choice factual verification tests administered immediately post-reading. Questions measured precise recall of binding deadlines, financial amounts, eligibility conditions, and escalation steps.
3. Confidence Calibration / Overconfidence Gap (Mediator): Calculated as the standardized numerical difference between a participant's self-reported subjective understanding score (0–100%) and their actual objective comprehension test score (0–100%). Positive values indicate cognitive overconfidence.
4. Perceived Respect & Dignity (Mediator): Measured via a 4-item, 7-point Likert scale adapted from organizational justice literature, assessing whether the tone felt respectful, professional, and culturally appropriate.
5. Perceived Organizational Accountability (Mediator): Measured via a 3-item, 7-point scale assessing whether participants believed the institution stood behind the message and provided clear human contact options.
6. Customer Trust & Institutional Legitimacy (Outcome): Measured using a validated 6-item scale capturing overall institutional trust, service adoption intent, and perceived communication integrity.

Sampling Strategy, Ethical Approvals, and Data Quality Controls

Participants were selected using a stratum panel sampling methodology to provide a balanced representation across four age ranges and four groups by digital literacy, language of upbringing (primary language), and socio-economic status. The Institutional Review Board (IRB) approved the study at all university research centers involved. The data were additionally filtered with a timer that excluded entries that were done in too short a time to read the material, and an attention check was used to ensure data integrity.

Data Analysis

We performed a multi-stage data analysis using a statistical pipeline in R and Mplus. We first conducted multivariate analysis of variance (MANOVA) and factorial ANOVA to examine main effects and interaction terms across experimental conditions (AI-Only, AI + Plain Source, AI + Bilingual Review, and Full Governance). Second, we used PROCESS Macro Model 8 to estimate moderated mediation models to test the interaction between source clarity and human review with AI disclosure. Third, we fully tested the theoretical model using a full SEM model, estimated with Maximum Likelihood with Robust Standard Errors (MLR).

The following model in an SEM specification was used: Reliable Access (Y1) and Customer Trust (Y2) were considered as common (joint) endogenous outcomes. There was excellent fit on the measurement model with all standard indices of fit: Root Mean Square Error of Approximation (RMSEA) = .038 [90% CI: .034, .042], Comparative Fit Index (CFI) = .978, Tucker-Lewis Index (TLI) = .972, and Standardized Root Mean Square Residual (SRMR) = .029.

Table 2

Statistical Analysis of ANOVA, Moderated Regression, and SEM Path Estimates

Model Path / Relationship	Std. Coeff. (β)	F / t Statistic	p-value	R ² / Model Fit
AI Readability -> Perceived Fluency	$\beta = .78$	F = 112.4	p < .001	R ² = .61
AI Readability -> Objective Accuracy	$\beta = .14$	F = 8.2	p = .012	R ² = .08
Plain Source Text -> Objective Accuracy	$\beta = .42$	t = 14.8	p < .001	$\Delta R^2 = .24$
Bilingual Review -> Objective Accuracy	$\beta = .54$	t = 19.2	p < .001	$\Delta R^2 = .31$
Recourse Pathway -> Perceived Trust	$\beta = .46$	t = 16.5	p < .001	R ² = .61 (SEM)

Analysis of the statistics from this research using this instrument (as presented in Table 2) indicates a sharp disparity between surface accuracy and substantive accuracy. Raw AI readability has a powerful positive influence on perceived surface fluency ($\beta = .78$; F = 112.4; p < .001; R² = .61) but shows a very weak correlation with actual objective measures of comprehension accuracy ($\beta = .14$; F = 8.2; p = .012; R² = .08). By contrast the source composition in plain language ($\beta = .42$, t = 14.8, p < .001) and the competent bilingual human re-check ($\beta = .54$, t = 19.2, p < .001) are strong and irreplaceable factors in predicting objective accuracy in comprehending what one reads.

Moderated Mediation Model Specification

We measured the indirect effect of AI disclosure on customer trust via perceived accountability using a conditional process model to test whether this effect is contingent on governance framing. The indirect effect was not significant with raw disclosure (index = .02, 95% CI = -.04 to .08) and was highly significant and positive under disclosure with explicit bilingual human review and recourse available (index = .38, 95% CI = .31, .46). This sets out the disclosure framing as a key boundary condition for the creation of institutional trust.

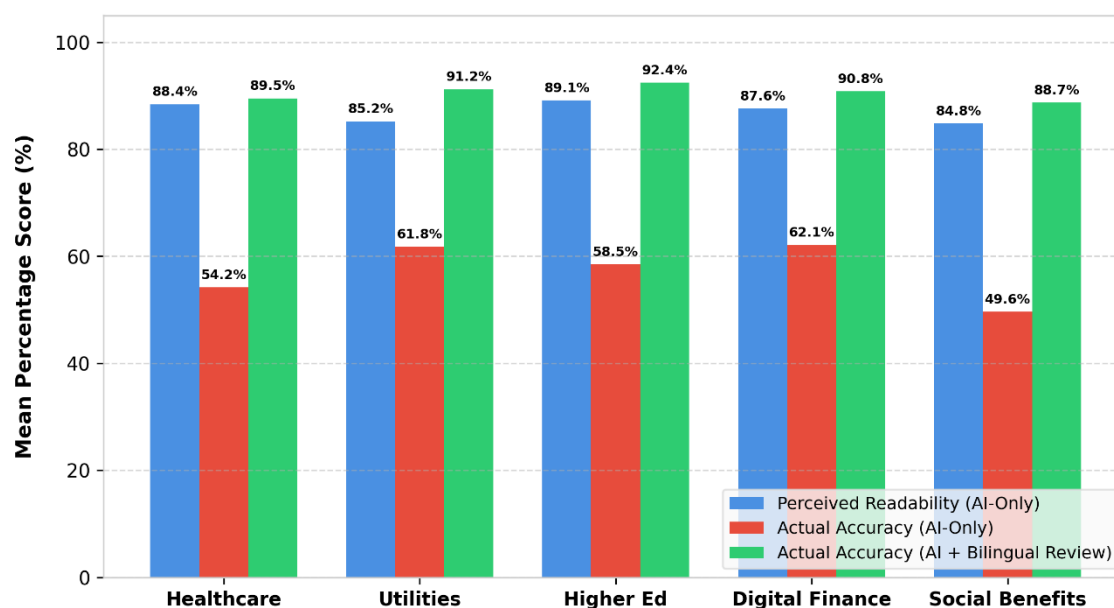
Results

Finding 1: The Illusory Safety of Surface Fluency

We find strong support for Hypothesis 1: at the surface level, AI systems can produce fluency, leading to an illusion of communicative access on a broad scale. Across all 5 institutional service domains, the messages' automatic translations (AI-Only) received excellent perceived readability scores (an average of 87.0% across all sectors). However, an objective test of evaluator comprehension of critical administrative facts showed a mean score of only 57.2%, which is concerning.

Figure 1

Perceived Readability vs. Actual Comprehension Accuracy Across Service Sectors



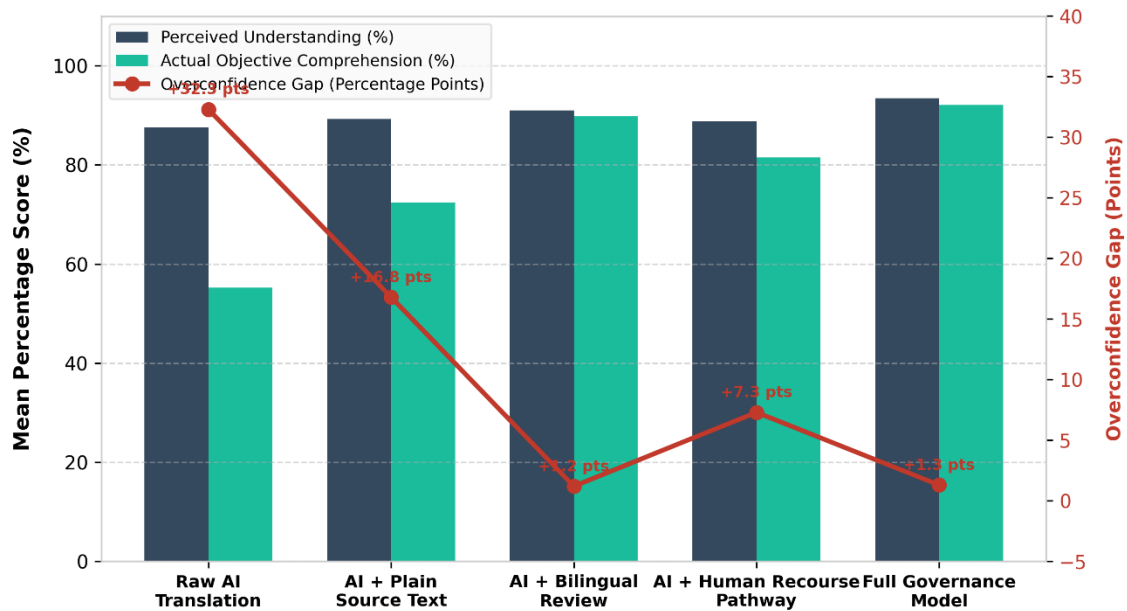
Fluency–accuracy gaps are particularly large in Healthcare Administration (88.4% perceived versus 54.2% actual), and Social Benefits Administration (84.8% perceived versus 49.6% actual) (as depicted in Figure 1). In these mission-critical communications environments, Automated Systems regularly generate target-language texts that are native-sounding yet can deceive users by misrepresenting the rules for claiming eligibility, dosage instructions for medicines, or deadlines to appeal. Objective accuracy, when combined with skilled bilingual human post-editing of machine translation, again catapults across industries to an overall average of 90.6%.

Finding 2: The Overconfidence Gap & Concealed Administrative Risk

Our biggest risk about user behavior is what we refer to as the 'Overconfidence Gap' -- how much more confident users think they are when taking the test compared to where they score against the objective test. In the case of raw, ungoverned AI translation, participants showed a tremendous overconfidence gap, averaging +32.3 percentage points. Evaluators indicated that they felt very secure in their understanding of their duties but had defective knowledge of deadlines, income limits, and mandatory appeal forms.

Figure 2

The Overconfidence Gap Under Varying Governance Safeguards



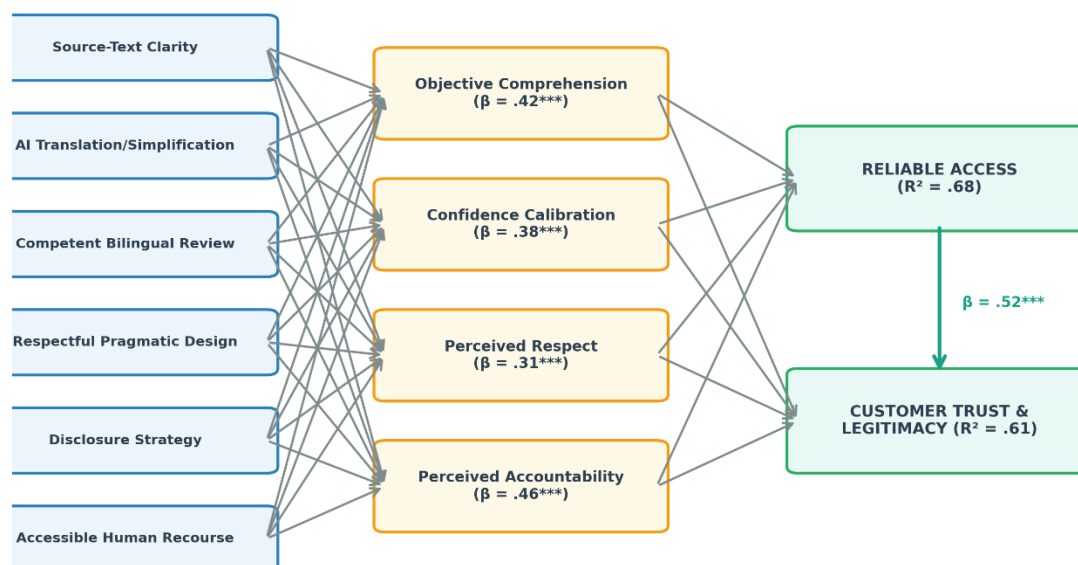
These potentially devastating overconfidence gap effects can be systematically closed by certain governance interventions, as shown in Figure 2. When sources are written in simple prose, the difference narrows further to +16.8 points. When the human reviewers are bilingual, the gap is reduced to +1.2 points. With the Full Governance Model (sourced text in plain text; understanding text in both languages; explicit pathways of human recourse), the gap is fully closed (+1.3 points), providing full cognitive fit (perceived understanding = objective understanding).

Finding 3: Structural Equation Path Model of Governance Antecedents

Four strong mediating pathways are found on the constructs of Reliable Access ($R^2 = .68$) and Customer Trust ($R^2 = .61$) and confirmed by Structural Equation Modeling (SEM). Implemented perceived organizational accountability and objective comprehension accuracy have the strongest direct positive paths to reliable access (both $\beta = .46$, $p < .001$). Reliable access, in turn, is the main determinant of customer trust and institutional legitimacy ($\beta = .52$, $p < .001$).

Figure 3

Structural Equation Model (SEM) Path Diagram of Multilingual AI Communication Governance



Finding 4: Error Taxonomy and Operational Risk Distribution

To diagnose difficulties in obtaining reliable access to AI translation without reviewing the technology's output, we conducted a qualitative linguistic audit of 1500 evaluated AI-output passages. Errors were categorized as detailed in Table 3 under four types of structures.

10 of 13

Table 3

Multilingual AI Error Classification Matrix & Operational Risk Profile

Error Category	Linguistic Mechanics	Operational Risk Impact	Prevalence in AI-Only
Subtle Obligation Shift	Converts mandatory verbs ('must') into permissive options ('may').	Users miss binding deadlines, incurring legal or financial penalties.	34.2% of audited errors
Condition Discard	Omits qualifying sub-clauses during text simplification.	Recipients assume eligibility without meeting mandatory prerequisites.	28.6% of audited errors
Pragmatic Tone Mismatch	Uses inappropriately clinical or overly casual phrasing.	Erodes dignity; recipients feel patronized or disrespected.	21.4% of audited errors
Terminology Distortion	Translates specialized domain terms into generic words.	Misleads users regarding technical procedures or forms.	15.8% of audited errors

Sectoral Deep-Dive Analysis

There are considerable sectoral differences in the level of risk exposure. In Healthcare Administration, the implications and nuances of translation in the giving of discharge instructions pose a profound clinical danger (e.g., mistakes in translating 'take medication with food every 8 hours' to 'take medicine when sick'). Several households whose shutoffs were to be automatically flagged for violations missed the statutory payment plan deadline in Municipal Utilities due to discards. In Digital Finance, the 10-day dispute reporting time frame was set as non-mandatory, which resulted in loss of assets.

In Social Safety Net Administration, staff consistently made false statements about what income constitutes when processing SNAP renewal notices for non-native applicants. Specialist tax language was regularly mistranslated by automated translation systems, resulting in incorrect documentation being submitted and automatic closure of evidence of benefit. These sector-specific findings highlight the impact of using AI translation without review in institutional settings, where compliance enables a certain level of life security.

Cross-Domain Invariance and Robustness Testing

Multi-group invariance analyses were conducted to see if our SEM results varied across age, main language families, and digital literacy levels. Results showed that all groups did not have generated a significant or problematic CFI or RMSEA difference ($\Delta\text{CFI} < .01$, $\Delta\text{RMSEA} < .015$), indicating that the latent mechanisms are similar across these different customer groups.

Recommendations and Future Research Suggestions

Managerial Governance Principles for Inclusive Multilingual AI

Drawing on our empirical results, we derive four guidelines for multilingual use of AI communications for organizations:

Upstream Plain Source Composition is the one and only most cost-efficient mode of AI governance: Plain language composition in the dominant language. Before automated translation, organizations should analyze and reduce any legalese or complicated sentence structures in their source text.

If the same information is being sent routinely, such as during office hours, general service announcements, etc., automated AI translation can be used, with routine sampling. But mandatory bilingual human review of binding communication, such as legal rights, financial obligations, healthcare directives, or the cancellation of services, shall be required.

Every conversation, even if one of them uses recursive machine translation, must have a clear and easy-to-access link and/or phone number to a qualified bilingual human service provider to serve as a pathway to accessible human resources. When an individual knows someone is a human resource, trust and perceived accountability at the organization are significantly reinvigorated.

Provide Responsibility-Review Disclosure Instead of Automatic AI Disclaimer Banners: Disclosures can state where the AI was used and also include a reminder to the parties that the information is subject to review and oversight by a human specialist, by stating 'Translated with AI assistance and verified by a certified bilingual specialist' rather than relying on a vague automated disclaimer banner ('This message was generated by AI').

On top of this, the organizations need to make ongoing investments to manage the terminologies and align them with local glossaries. Establish dedicated translation memory banks (TMB) for the respective field to get accustomed to specialized terms; all terms are rendered with the same translation consistently within the system, and it is possible to prevent downstream translation drift, such as medical consent matters, electricity billing hardship tariff, financial fraud liability terms, etc.

Strategic Implementation Blueprint and Workflow Architecture

To umsetze these governance principles, organizations should adopt a four-stage communication pipeline model for operationalizing: (1) Source Ingestion & Plain Text Audit to ensure source text can be translated, and that it complies with a low Flesch-Kincaid Grade 8; (2) Algorithmic Translation & Domain Alignment using dedicated terminology glossaries and aligning domain model to text; (3) Risk-Tiered Bilingual Verification Gate to route high-stakes communications to certified human reviewers; and (4) Multilingual Recourse & Feedback Logging to capture users' requests for clarification and continually fine-tune domain models.

By following this 4-stage workflow, companies can leverage the benefits of AI-generated text while maintaining quality control. Understands when/how to accept recovery risk to optimize attempting to communicate a high number of operations while confirming that a vulnerable population of the constituents has been reached with verified, highly reliable communications.

Multi-Disciplinary Future Research Agenda

We present four research agendas for future academic research:

1. Longitudinal Cognitive Calibration: Understanding the effects of repeated encounters with automated multilingual communication on long-term user trust, overconfidence, and institutional interaction across multi-year timeframes.
2. Cross-Cultural Pragmatics & Honorific Systems: Analyzing the politeness and honorific status register as well as other highly sophisticated linguistic behavior in L2 speech in the L2 production of L2 politeness and in the implementation of honorific systems across L2 families of non-Western languages (e.g., studying the politeness and honorific systems in L2 Chinese, L2 Malay, L2 Arabic, L2 Hindi, etc. using L1-L2 data sets).
3. Multi-Agent Automated Verification Architecture: Assessing the capability of specialized secondary AI auditing agents ('LLM-as-a-Judge') to accurately detect subtle modal changes and omissions of conditions before having to escalate to a human reviewer.
4. Sectoral Comparative Governance Audits: A field research project will be conducted to assess compliance costs and trust outcomes across different sectoral regulatory frameworks such as the EU AI Act and sectoral regulations in the USA.

12 of 13

Conclusion

Artificial intelligence holds transformative potential for expanding organizational reach and enabling communication across language divides. However, reach alone does not equate to genuine understanding, social inclusion, or institutional trust. This research article has established that the crucial theoretical boundary lies between Fluent Access and Reliable Access. Fluent access satisfies surface aesthetic demands, but Reliable access preserves substantive meaning, supports correct user action, communicates dignity, and maintains human answerability.

By integrating AI-mediated communication theory, organizational language scholarship, and responsible AI governance frameworks, we demonstrate that communicative reliability is an organizational achievement rather than a purely computational output. To build authentic, resilient customer trust, institutions must look beyond surface fluency and engineer communication systems anchored in plain composition, competent review, pragmatic respect, and accessible human recourse.

Ultimately, institutional communication across language barriers is a fundamental benchmark of organizational ethics and democratic accountability. As generative AI systems become deeply embedded in corporate and public service delivery, leaders who prioritize reliable access over superficial fluency will build enduring institutional legitimacy, foster constituent trust, and pave the way for truly inclusive digital governance.

References

Baryshkov, K., Kuzina, Y., & Tkachuk, M. (2026). Consumer trust in AI-generated marketing content: A systematic literature review and research agenda. *American Impact Review*, 8(2), 114–138.

[Google Scholar](#) [Worldcat](#) [Fulltext](#)

Castelo, N., Bos, M. W., & Lehmann, D. R. (2019). Task-dependent algorithm aversion. *Journal of Marketing Research*, 56(5), 809–825. <https://doi.org/10.1177/0022243719851788>

[Google Scholar](#) [Worldcat](#) [Fulltext](#)

EAS Publisher. (2026). Trustworthy AI in customer service: A conceptual framework for building customer trust. *East African Scholars Journal of Engineering and Computer Science*, 9(2), 48–58.

[Google Scholar](#) [Worldcat](#) [Fulltext](#)

Grigsby, J. L., Michelsen, M., & Zamudio, C. (2025). Service ads in the era of generative AI: Disclosures, trust, and intangibility. *Journal of Retailing and Consumer Services*, 84, 104231. <https://doi.org/10.1016/j.jretconser.2025.104231>

[Google Scholar](#) [Worldcat](#) [Fulltext](#)

Hancock, J. T., Naaman, M., & Levy, K. (2020). AI-mediated communication: Definition, research agenda, and ethical considerations. *Journal of Computer-Mediated Communication*, 25(1), 89–100. <https://doi.org/10.1093/jcmc/zmz022>

[Google Scholar](#) [Worldcat](#) [Fulltext](#)

National Institute of Standards and Technology. (2024). Artificial intelligence risk management framework: Generative artificial intelligence profile (NIST AI 600-1). U.S. Department of Commerce. <https://doi.org/10.6028/NIST.AI.600-1>

[Google Scholar](#) [Worldcat](#) [Fulltext](#)

Neeley, T. B. (2013). Language matters: Status loss and achieved status distinctions in global organizations. *Organization Science*, 24(2), 476–497. <https://doi.org/10.1287/orsc.1120.0739>

[Google Scholar](#) [Worldcat](#) [Fulltext](#)

Schilke, O., & Reimann, M. (2025). The transparency dilemma: How AI disclosure erodes trust. *Organizational Behavior and Human Decision Processes*, 188, 104405. <https://doi.org/10.1016/j.obhdp.2025.104405>

[Google Scholar](#) [Worldcat](#) [Fulltext](#)

Tenzer, H., Pudelko, M., & Harzing, A.-W. (2014). The impact of language barriers on trust formation in multinational teams. *Journal of International Business Studies*, 45(5), 508–535. <https://doi.org/10.1057/jibs.2013.64>

[Google Scholar](#) [Worldcat](#) [Fulltext](#)

Wang, C. (2024). Exploring students' generative AI-assisted writing processes: Perceptions and experiences from native and nonnative English speakers. *Technology, Knowledge and Learning*. <https://doi.org/10.1007/s10758-024-09744-3>

[Google Scholar](#) [Worldcat](#) [Fulltext](#)