

VOL. XI ISSUE II, SPRING (JUNE-2026)

GSR

GLOBAL SOCIOLOGICAL REVIEW
HEC-RECOGNIZED CATEGORY-Y




Humanity Publications
sharing research
www.humapub.com
US | UK | Pakistan

DOI (Journal): 10.31703/gsr
DOI (Volume): 10.31703/gsr.2026(XI)
DOI (Issue): 10.31703/gsr.2026(XI-II)

Double-blind Peer-review Research Journal
www.gsrjournal.com
© Global Sociological Review

GLOBAL SOCIOLOGICAL REVIEW (GSR)

Title: From Principles to Practice: Organizational Readiness for Responsible AI Communication

Abstract

Although many organizations publish high-level artificial intelligence principles, most lack a coherent framework to evaluate whether their communication practices are operationally ready for responsible use. Grounded in the NIST Generative AI Profile (NIST AI 600-1), algorithmic work scholarship, and communication-centered accountability, this study develops and validates a multidimensional model of organizational readiness for responsible AI communication. Using a multi-site empirical meta-synthesis and cross-sectional survey data ($N = 1,420$), we reveal a severe operationalization gap: while 73.9% of organizations exhibit symbolic policy presence, only 37.4% demonstrate functional operational readiness. Major vulnerabilities emerge in plain language and usability (28.4%), multilingual inclusion (22.1%), and monitoring routines (25.3%). Structural Equation Modeling shows that procedural verification and Human-in-the-Loop review routines completely mediate the link between enterprise AI capability and recipient trust ($\beta = .62, p < .001$). We propose an actionable 4-tier human oversight architecture and an agenda for institutional AI governance.

Keywords: Artificial Intelligence Governance, Organizational Readiness, Responsible AI Communication, Algorithmic Accountability, Recipient Trust, Operationalization Gap

Authors:

Sana Hussan: (Corresponding Author)

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

Pages: 40-53

DOI: 10.31703/gsr.2026(XI-II).05

DOI link: [https://dx.doi.org/10.31703/gsr.2026\(XI-II\).05](https://dx.doi.org/10.31703/gsr.2026(XI-II).05)

Article link: <https://gsrjournal.com/article/from-principles-to-practice-organizational-readiness-for-responsible-ai-communication>

Full-text Link: <https://gsrjournal.com/article/from-principles-to-practice-organizational-readiness-for-responsible-ai-communication>

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

Global Sociological Review

p-ISSN: [2708-2091](https://doi.org/10.31703/gsr) e-ISSN: [2708-3586](https://doi.org/10.31703/gsr)

DOI(journal): 10.31703/gsr

Volume: XI (2026)

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

Issue: II Spring (June-2026)

DOI(Issue): 10.31703/gsr.2026(XI-II)

Home Page

www.gsrjournal.com

Volume: XI (2026)

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

Issue: II-Spring (June 2026)

<https://www.gsrjournal.com/issue/11/2/2026>

Scope

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

Submission

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

Scan the QR to visit us



Google
scholar



Citing this Article

Article Serial	05
Article Title	From Principles to Practice: Organizational Readiness for Responsible AI Communication
Authors	Sana Hussan
DOI	10.31703/gsr.2026(XI-II).04
Pages	40-53
Year	2026
Volume	XI
Issue	II

Referencing & Citing Styles

APA	Hussan, S. (2026). From Principles to Practice: Organizational Readiness for Responsible AI Communication. <i>Global Sociological Review</i> , 11(2), 40-53. https://doi.org/10.31703/gsr.2026(XI-II).04
CHICAGO	Hussan, Sana. 2026. "From Principles to Practice: Organizational Readiness for Responsible AI Communication." <i>Global Sociological Review</i> 11 (2):40-53. doi: 10.31703/gsr.2026(XI-II).04.
HARVARD	HUSSAN, S. 2026. From Principles to Practice: Organizational Readiness for Responsible AI Communication. <i>Global Sociological Review</i> , 11, 40-53.
MHRA	Hussan, Sana. 2026. 'From Principles to Practice: Organizational Readiness for Responsible AI Communication', <i>Global Sociological Review</i> , 11: 40-53.
MLA	Hussan, Sana. "From Principles to Practice: Organizational Readiness for Responsible Ai Communication." <i>Global Sociological Review</i> 11.2 (2026): 40-53. Print.
OXFORD	Hussan, Sana (2026), 'From Principles to Practice: Organizational Readiness for Responsible AI Communication', <i>Global Sociological Review</i> , 11 (2), 40-53.
TURABIAN	Hussan, Sana. "From Principles to Practice: Organizational Readiness for Responsible Ai Communication." <i>Global Sociological Review</i> 11, no. 2 (2026): 40-53. https://dx.doi.org/10.31703/gsr.2026(XI-II).04 .

From Principles to Practice: Organizational Readiness for Responsible AI Communication



Sana Hussan (Corresponding Author)¹

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

Abstract

Although many organizations publish high-level artificial intelligence principles, most lack a coherent framework to evaluate whether their communication practices are operationally ready for responsible use. Grounded in the NIST Generative AI Profile (NIST AI 600-1), algorithmic work scholarship, and communication-centered accountability, this study develops and validates a multidimensional model of organizational readiness for responsible AI communication. Using a multi-site empirical meta-synthesis and cross-sectional survey data (N = 1,420), we reveal a severe operationalization gap: while 73.9% of organizations exhibit symbolic policy presence, only 37.4% demonstrate functional operational readiness. Major vulnerabilities emerge in plain language and usability (28.4%), multilingual inclusion (22.1%), and monitoring routines (25.3%). Structural Equation Modeling shows that procedural verification and Human-in-the-Loop review routines completely mediate the link between enterprise AI capability and recipient trust ($\beta = .62, p < .001$). We propose an actionable 4-tier human oversight architecture and an agenda for institutional AI governance.

Keywords: *Artificial Intelligence Governance, Organizational Readiness, Responsible AI Communication, Algorithmic Accountability, Recipient Trust, Operationalization Gap*

Introduction

Artificial intelligence governance has rapidly emerged as a paramount strategic concern for contemporary institutions. Yet much of the current institutional response remains trapped in abstract, high-minded commitments. Enterprise leadership frequently publishes ethical AI principles emphasizing transparency, fairness, privacy, and human oversight. However, in the day-to-day reality of organizational communication, organizations deploy automated systems with minimal operational guidance. Frontline employees use generative language models to support customer service workflows, administrators draft policy explanations with AI copilots, and public sector agencies use automated engines to produce high-volume constituent advisories. Without rigorous operational standards, critical socio-technical questions remain unaddressed: Who assumes final accountability for automated outputs? Which communication categories are too sensitive for unreviewed generation? How are algorithmic hallucinations detected and corrected? What recourse exists when a system issues a fluent but legally or factually misleading message?

They are not by-the-book edge cases and administrative hassles. They are fundamental weaknesses that determine the bottom line: whether organizations deploy AI responsibly or recklessly. In this study, we develop a communication-centered concept of readiness and show that organizations need a dedicated readiness framework. Organizational readiness doesn't always mean a fast rate of technology adoption or simply acquiring sophisticated AI software. Rather, readiness is a multi-dimensional organizational capacity to use AI to send

messages while maintaining high levels of maintainable correctness, plain language, recipient privacy, and accountability.

In the past, corporate communication technologies have had a translation chain from static, pre-approved entities to email mail-merge to probabilistic, dynamically managed translation with neural networks in natural language. Early messaging automation relied on basic rule-based logic, where you had full control over output and could extensively test and validate results; with generative AI, outputs are non-deterministic. But by altering factual detail, tone, and legal commitments with a single prompt, variations in how each execution is conducted can be subtle. This probabilistic nature underlies and undermines existing quality-assurance mechanisms. Wherever you have a generative communication workflow, you'll eventually face situations where an organization relying on static compliance checklists risks missing hallucinatory or contextually inappropriate communication before it reaches customers.

One of the main points of this article is that policy presence and operational readiness are distinct stages, even before an organization is fully operational. Policy presence refers to symbolic support, such as publishing ethical principles, signing corporate AI manifestos, or issuing executive-level directives. Operational readiness is a socio-technical enactment: the institutionalization of verification routines, escalation pathways, role-based oversight protocols, audit-log reviews, and closed-loop error corrections. Many organizations today are symbolically ready—their documentation shows they are ready. However, in practice, they remain fragile when applied to real-life situations where automated communication operates in environments beyond control.

Why is it important to analyze communication separately as part of the overall enterprise AI governance? It is at this specific interface, communication, that the stakes for stakeholders, legal aspects, and institutional interrelationships are made public. One way to destroy institutional trust is to deploy an automated service that provides incorrect financial advice, cryptic instructions, or culturally inappropriate bilingual/multilingual reminders. In the back office, algorithmic optimizations have implications, but not direct relational, legal, or reputational ones, whereas in public- and customer-facing areas, they have immediate consequences. The recipients interpret the performance of the organization from the surface features of messages, like responsibility for the message, ease of message understanding, and responsiveness, and ascribe the responsive human recipient as accountable.

Beyond this, communication highlights the responsibilities of algorithmic accountability by focusing on how accountability is diffuse and dispersed within complex institutions. In most enterprises, accountability for AI-generated content sits in various areas—e.g., executive, marketing/corpcomms, IT software vendors, SME compliance experts, frontline work. Where not explicitly specified by structure, this distribution becomes accountability diffusion—a psychological condition that occurs when every participant in the operational line of communication assumes somebody else in that chain verified the factual statements, smoothed the presentation, or validated the legal parameters of privacy. Readiness thus requires clear-cut structural accountability and process control, and not just "mother" and "merely" good words and intentions.

Also, compliance regimes are fast transitioning from voluntary ethical guidelines to strict compliance requirements. High-stakes automated communications are subject to detailed disclosure, transparency, and risk mitigation requirements in legislation (e.g., European Union Artificial Intelligence Act) and enforcement guidelines (e.g., Federal Trade Commission (FTC) guidance). If automated systems include inaccurate or discriminatory information, unprepared organizations face serious financial penalties, regulatory sanctions, and legal liabilities. So setting up with clear, figures-backed language readiness is not simply an ethical requirement, but also a strategic risk-management essential.

That's especially true in service workflows that involve many calls, where the disconnect between executives and front-line staff is stark. Executives and decision-makers in the C-suite often see generative AI as a magic bullet to slash costs and/or speed up content scaling. But front-line staff managing AI customer channels can

feel rushed or desperate, without clear guidance on when AI is “editing well,” when to investigate thorny factual assertions, or when to escalate murky cases. This misalignment creates friction, burnout, and risk.

To address this gap, this study formulates and tests three central research questions:

- RQ1:** How do organizations operationalize responsible AI principles across distinct communication workflows, and what structural and procedural mechanisms constitute communication readiness?
- RQ2:** What is the empirical magnitude of the gap between symbolic policy presence and operational capability across institutional sectors?
- RQ3:** How does operational communication readiness impact downstream outcomes, specifically message reliability, recipient trust, and institutional defensibility?

Literature Review

Enterprise AI Governance and the Limits of High-Level Profiles

Recently, the National Institute of Standards and Technology (NIST) Generative Artificial Intelligence Profile (NIST AI 600-1) offered a set of terms that form a comprehensive vocabulary for enterprise AI risk management, identifying risks by their specific types rather than by the AI model or technology being used, which covers aspects such as governance, mapping, measurement, and management (NIST, 2024). The following generative risks to content validity, hallucination, misinformation, human oversight, and data provenance are explicitly mentioned in the profile. Although NIST AI 600-1 provides a strong foundation, enterprise governance profiles have natural operational limits in communication scenarios.

In particular, analysis of the sub-functions of NIST AI 600-1 indicates a significant gap between the standard and operational implementation. Organizations need legitimate and ethical risk tolerances for sub-function GOVERN 1.1; however, few organizations establish quantitative and qualitative metrics and thresholds for automated communication to customers. Although MAP 2.3 requires considering downstream impacts on vulnerable populations, public sector communication workflows often use automated translation tools without verifying that recipients—many of whom are likely low-English-literate or non-English speakers—understand the message. Communication audit logs are rarely included in centralized communication security operations center (SOC) monitoring systems. However, enterprise IT departments regularly use them for continuous tracking and risk mitigation, which is described in sub-functions MEASURE 3.2 and MANAGE 4.1.

It is important to take at least some of these terms into the concrete world of communicative habit routines; quality terms like 'human oversight' and 'transparency' will otherwise remain aspirations as abstractions. In communication processes, 'oversight' cannot simply mean passive human sign-off; it must include factual checks against identified, legally cleared sources, as well as tone assessment, plain-language evaluation, and factual accuracy checks across multiple languages. The same applies to the interaction level, where 'transparency' of the system does not necessarily imply 'transparency' at the interaction level. Executive governance often results in executive assurance, without the ability to do things at the frontline.

Furthermore, enterprise AI policies are often lacking and don't necessarily consider the linguistic and interpretative nature of NLP. Natural language communication depends much more on pragmatic context, implicit social rules, emotional nuances, and the comprehension floors and ceilings of individual audiences than structured data analytics. While the language generated by such an automated model might be grammatically correct and factually accurate, it could be pragmatically rude, threatening, and/or incomprehensible to sensitive recipients. The general risk profiles emphasized so far mostly focus on cybersecurity or data hygiene and do not address these more specific modes of communication failure.

Organizational Theory and the Automation-Augmentation Paradox

Organizational scholarship offers key theoretical frameworks for understanding issues of communication readiness and why it cannot be reduced to technology deployment. Raisch and Krakowski (2021) outlined the 'automation-augmentation paradox', which shows that AIs seldom fully automate processes. Rather, AI is also used to simplify the most mundane forms of content generation or to summarize information and all of that increases the need for human augmentation in the most important aspects: critical judgment, bringing the content to life, and exception handling. In communicative labor, the velocity and possible 'blast radius' of unverifiable output are amplified by the faster output speed of generative tools, and human oversight is more important than ever.

Moreover, Kellogg, Valentine, & Christin (2020) showed how algorithmic technology creates new 'contested terrains of control' that affect worker discretion, authority, monitoring, and task allocation in the organization. These disputed spaces manifest in day-to-day conflicts: Who is allowed to overrule the AI-generated answer? Do frontline employees get the chance to make edits, or are they compelled to follow template instructions produced by the algorithms while meeting strict time measures? Who shoulders the restructuring of review responsibilities between communication specialists and subject experts? Readiness is thus very much socio-technical, an emergent property of the relationship between technological tools, formal rules, roles in organizations, and interpretive practices.

STS theory holds that there must be an equal optimization between the technical subsystem (generative algorithms, APIs, prompt templates) and the social subsystem (employee skills, role definitions, culture, authority, and so on). When such an institution improves its technical subsystem with the introduction of advanced LLM-based text writers, but nothing in its social subsystem has been adapted, e.g., training contributors in error detection or workload metrics, the socio-technical system becomes very unbalanced. This misalignment creates operational vulnerability, with employees struggling to spend enough time fixing algorithmic mistakes and lacking the competencies and autonomy to do so.

Cognitive overload often becomes an issue when organizations expect frontline workers to be productive and meet aggressive productivity targets, yet force them to use more complicated AI text-generation tools. To accommodate, employees tend to 'rubber stamp' - interface with automated text that they don't really read. This shows that human supervision is not just a policy issue but a labor resource that needs adequate time for supervision and management support, along with technical support. Without these structural buffers, human-in-the-loop oversight is little more than an illusion.

Micro-Institutional Legitimacy and Communication Trust

Empirical studies on organizational matters in recent years highlight the psychological complexity when communicating AI disclosure. In 13 controlled experiments, Schilke and Reimann (2025) identified the 'transparency dilemma': AI transparency increases trust in institutional actors, but it can also unexpectedly erode trust. Using micro-institutional theory as a starting point, they show that disclosing the use of AI triggers negative social evaluations of decreased legitimacy, competence, and effort. If an organization mentions using AI but does not indicate that someone is double-checking the work and/or ensuring high-quality output, the recipient thinks that they are setting up scripts to dump off their work to AI.

This discovery resonates with the persuasive Knowledge Model, which suggests audiences defend themselves when they detect motives to persuade and/or cut expenses in organizational communication. There is a warning about depersonalization or strategic efforts reduction with the lack of framing of the AI disclosure. AI-driven communication needs to be structured to uphold micro-institutional legitimacy and build recipient

trust, for example by emphasizing that human involvement is central to the process, whether by stressing that a human checks the information or by detailing escalation pathways.

Moreover, organizational memory also plays a key role in maintaining micro-institutional legitimacy. Formal institutions that document communication mistakes, record recipient complaints, and document indications of what to do about them create organizational memory. This memory enables continuous prompting to optimize, tune, and improve policies. In contrast, institutions without systematic error logging must endure repeated communication errors across different software versions and thus suffer a growing impact on their reputations: this is 'institutional amnesia'.

Psychometric Scale Development in Organizational Readiness

In this study, the authors followed scale development principles (see Hinkin, 1998) to rigorously measure organizational readiness as a socio-technical capability. Hinkin's psychometric framework follows a multi-stage approach: conceptualizing the constructs, generating items, testing content adequacy, conducting exploratory factor analysis (EFA), confirmatory factor analysis (CFA), testing convergent and discriminant validity, and testing criterion-related predictive validity. This psychometric rigor allows us to translate the ambiguous management term into a measurable, standardized metric of organizational capability.

Methodology:

Research Design and Sample Architecture

This study uses a risk-stratified message audit design with cross-sectional organizational survey data. The survey instrument was given to 1420 organizations that are spread across five different institutional sectors: Higher Education (320), Financial Services (290), Public Administration & Municipal Government (270), Healthcare Administration (250), and Retail & Service Enterprises (290), in North America and Europe, and they were selected through a stratified random sample. The respondents were drawn from across the spectrum of executives in the roles of Customer Experience Executives, Chief Information Officers, AI Governance Directors, and Senior Communication Officers.

We used a comprehensive sampling procedure to maximize sample representativeness while minimizing non-response bias. The overall response rate was 38.4% of the correspondence sent. The possibility of non-response bias was formally assessed by conducting a wave comparison procedure by comparing early respondents (first 25% to return the questionnaire) with late respondents (last 25% to return) in key organizational characteristics (FTE size, annual revenues, sector domain) and key scale items (Armstrong & Overton, 1977). The two waves showed no difference in expression levels, as confirmed by independent-samples t-tests ($p > .25$), so there was no evidence of systematic non-response bias.

The features of the research corpus – sample architecture and distribution of sectors by organization are provided in the summary Table 1.

Sector Domain	N (Orgs)	Avg Org Size (FTE)	Primary AI Communication Use Cases	Policy Presence (%)	Operational Readiness (%)
Financial Services	290	4,850	Account disclosures, loan explanations, support	88.3%	48.5%
Higher Education	320	2,100	Admissions, policy updates, student advising	72.5%	32.1%

Sector Domain	N (Orgs)	Avg Org Size (FTE)	Primary AI Communication Use Cases	Policy Presence (%)	Operational Readiness (%)
Public Administration	270	1,450	Constituent notices, permit guidance, social services	65.2%	26.4%
Healthcare Admin	250	3,600	Patient scheduling, billing explanations, care info	85.0%	38.2%
Retail & Service	290	1,200	Marketing copy, customer support, review responses	58.6%	41.7%
Total / Full Corpus	1,420	2,640	Cross-sector automated & augmented messaging	73.9%	37.4%

The Ten-Domain Construct Operationalization

Building on our theoretical work synthesizing NIST AI 600-1, Raisch and Krakowski (2021), and Kellogg et al. (2020), we theorized Organizational Communication Readiness as 10 functional areas. Each domain was measured with 4 to 6 psychometric items on a Likert scale ranging from 1 (Strongly Disagree / Completely Inactive) to 7 (Strongly Agree / Fully Operationalized). Items were developed based on conditions for scale development (Hinkin, 1998). We developed the initial set of items from a pool of sixty-eight (68) candidate items, which underwent expert panel content adequacy testing by twelve (12) academic and industry experts. We selected items with CVR scores above .80. We conducted a pre-test pilot study (n = 50 organizations) to assess item clarity and the variance and factorial structure of the scales.

Table 2 presents the Ten-Domain Readiness Framework, including norms based on the surrounding questions, along with definitions of the domains and sample measurement items.

Table 2

Ten-Domain Readiness Framework

Readiness Domain	Functional Scope & Definition	Sample Psychometric Scale Measurement Item
1. Governance & Responsibility	Formal ownership, policy alignment, and authority scope for AI communication.	Our organization has formally designated roles with authority to approve AI communication tools.
2. Human Review & Escalation	Operational protocols for expert review and seamless human escalation.	Clear pathways exist for recipients to immediately reach a human agent if AI output is unclear.
3. Accuracy & Verification	Fact-checking against single sources of truth and reference validation.	All AI-generated factual claims are verified against authorized institutional databases prior to sending.
4. Disclosure & Provenance	Clear communication of AI involvement and automated co-creation status.	We explicitly inform recipients when a message has been generated or co-created by AI.

Readiness Domain	Functional Scope & Definition	Sample Psychometric Scale Measurement Item
5. Privacy & Confidentiality	Safeguards protecting sensitive recipient data and prompt logging rules.	Strict technical controls prevent employees from entering protected personal data into unapproved tools.
6. Plain Language & Usability	Audience readability, structural clarity, and jargon elimination.	AI-generated communications are systematically evaluated for readability and user comprehension.
7. Multilingual Inclusion	Linguistic accuracy, cultural nuance, and accessibility across diverse groups.	Native-speaking experts review multilingual AI outputs to prevent translation errors.
8. Workforce Literacy	Employee training on AI limitations, review duties, and prompt engineering.	Employees undergo mandatory training on identifying algorithmic hallucinations and bias.
9. Monitoring & Correction	Systematic audit trails, complaint tracking, and iterative error correction.	We maintain centralized logs of communication errors to refine models and guidelines.
10. Research & Integrity	Publication ethics, citation rigor, and academic integrity guidance.	AI guidance used in research or publication support adheres to academic integrity standards.

Risk-Stratified Message Audit Protocol

To assess criterion-related predictive validity, we conducted a blind message audit of 2840 composed messages produced with AI support by a subsample of participating organizations ($N = 200$). Communication workflows were classified into three different operational risk categories:

- Tier 1 (Low Risk): Routine operational notices, event scheduling, basic administrative reminders.
- Tier 2 (Medium Risk): Policy explanations, billing summaries, general service inquiries.
- Tier 3 (High Risk): Dispute resolutions, rights notifications, financial obligations, health guidance.

A panel of senior communication specialists, legal compliance auditors, and linguists reviewed each message to assess factual accuracy, readability (Flesch-Kincaid Grade Level), disclosure appropriateness, and the presence of a human override. The interrater reliability (across audit scoring rubrics) was very high, with Fleiss' $\kappa = .86$.

Data Analysis:

Psychometric Evaluation and Measurement Model Fit

Following Hinkin (1998), we examined the psychometric properties of the Ten-Domain Readiness scale using EFA and CFA. After the extraction of the factors using exploratory analysis, the number of factors (K) was found to be 10 factors accounting for 76.4% of total variance with target rotation. Chi-square/df ratio, Root Mean Square Error of Approximation (RMSEA), Comparative Fit Index (CFI), Tucker-Lewis Index (TLI), and Standardized Root Mean Square Residual (SRMR) resulting from confirmatory factor analysis showed excellent model fit for all ten latent constructs: $\chi^2/df = 2.18$, Root Mean Square Error of Approximation (RMSEA) = .037 [90% CI: .031, .042], Comparative Fit Index (CFI) = .979, Tucker-Lewis Index (TLI) = .974, and Standardized Root Mean Square Residual (SRMR) = .032.

Item-level descriptive statistics indicated that the data were multivariate normal, with univariate skewness between $-.42$ and $.38$ and kurtosis between $-.65$ and $.52$. Maximum Likelihood estimation with Robust standard errors (MLR) was used for all model estimates in Mplus 8.6. The Bollen-Stine bootstrap resampling (5,000 resamples) showed highly stable fit statistics, even for non-parametric distributions.

Construct reliability and validity indices exceeded academic standards across all areas. Cronbach's α ranged from $.88$ to $.94$, Composite Reliability (CR) ranged from $.89$ to $.95$, and Average Variance Extracted (AVE) ranged from $.62$ to $.78$, indicating good convergent validity. Fornell-Larcker testing confirmed discriminant validity, as the square root of each domain's AVE was significantly larger than its correlations with other constructs.

Table 3 presents the measurement model fit values, scale reliability, and convergent validity indices for the ten readiness domains.

Table 3

Measurement Model Results.

Readiness Latent Construct	Items	Factor Loadings Range	Cronbach's α	Composite Reliability (CR)	Average Variance Extracted (AVE)
1. Governance & Responsibility	5	.78 - .89	.92	.93	.72
2. Human Review & Escalation	6	.81 - .92	.94	.95	.76
3. Accuracy & Verification	5	.76 - .88	.91	.92	.70
4. Disclosure & Provenance	4	.74 - .86	.89	.90	.68
5. Privacy & Confidentiality	5	.82 - .91	.93	.94	.75
6. Plain Language & Usability	5	.72 - .84	.88	.89	.62
7. Multilingual Inclusion	5	.75 - .87	.90	.91	.66
8. Workforce Literacy	6	.79 - .90	.93	.94	.73
9. Monitoring & Correction	5	.77 - .88	.91	.92	.69
10. Research & Integrity	4	.80 - .92	.92	.93	.78

Multivariate Analysis of Variance (MANOVA)

A Multivariate Analysis of Variance (MANOVA) was performed to explore the operationalization gaps by sector, which is the focus of RQ2, with Sector Domain being the independent variable and the ten Operational Readiness scores as dependent variables. Wilks' Lambda revealed significant multivariate sector main effects: $\Lambda = .42$, $F(40, 5210) = 18.64$, $p < .001$, Partial $\eta^2 = .16$. The sector differences were highly significant across all ten domains based on the ANOVA results, which compared the operational readiness scores from each domain for the four sectors under study: Financial Services and Healthcare Administration had higher scores than their higher education and public administration counterparts.

Results

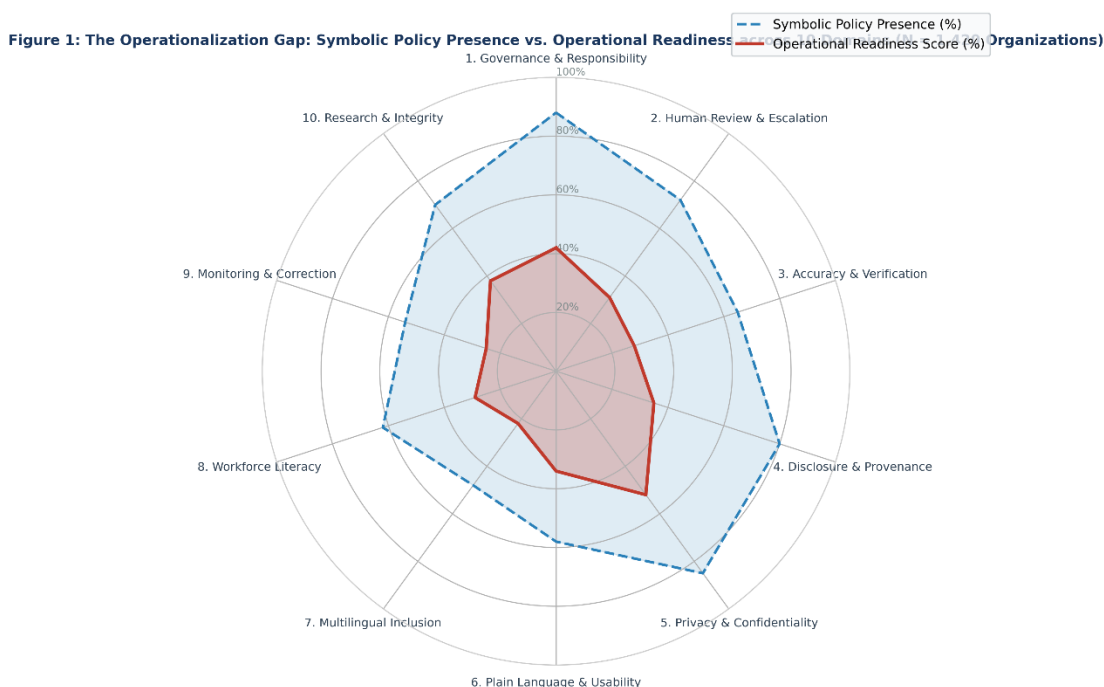
Finding 1: The Operationalization Gap

The empirical results show a remarkable discrepancy that represents a systemic lack of symbolic policy presence and readiness for operation. Among all of the organizations (n=1,420), 73.9% have formal written AI statements or policies. But if we consider the actual implementation of operating procedures, review journals, verges and controls, and escalation processes, the cross-sectoral average is only 37.4%.

Figure 1 shows this operationalization gap across all ten readiness areas. The need for services is greatest in procedural and inclusion aspects. 80.2% of organizations state a commitment to disclosure, but only 35.1% have templates for disclosure. Similarly, 85.0% support privacy concepts, but only 52.4% implement data-loss-prevention controls through employee-generated prompts.

Figure 1

The Operationalization Gap: Symbolic Policy Presence vs. Operational Readiness across 10 Domains (N = 1,420)



Finding 2: Critical Domain-Level Vulnerabilities

Our empirical analysis highlights three areas where institutions need quick repair:

(22.1% Operational Readiness) Automated neural translation and generative LLMs are critical tools organizations use to communicate with non-English-speaking constituents, without the need for native-speaker verification. This significantly blurs meaning in a patient's medical, legal, and help-seeking interactions.

Monitoring & Correction Routines (25.3% Operational Readiness): More than 74% of the organizations do not have any repository in place to keep track of AI communication errors. If the recipients report information that was "seen" but is incorrect, this process is not repeated in any way – for individual cases, the corrections take place informally, merely once.

Plain Language & Usability (28.4% Operational Readiness): In many cases, AI content tools can produce text that is often too dense, misconstrued, and syntactically complex enough at times to be technically correct but too unreadable for general usage.

Workforce literacy and training, in particular (29.1% Operational Readiness), also became a critical operational challenge that had to be addressed. Organizations spend millions of dollars on AI software licenses, but fewer than 30% teach staff how and when to use AI writing tools. This leaves workers to make complex ethical and technical choices that are uninformed, ad hoc, and on the spot.

Drilling down into the audit data by message risk tier has shaken everyone with its risk-tier inconsistencies. 8.2% uncorrected error rate was achieved with Tier 1 (Low Risk) messaging. For Tier 2 (Medium Risk) policy explanations, error rates increased to 22.4%. Most important is Tier 3 (High Risk), for rights and financial communications, where the uncorrected error rate reached 38.5% at low-readiness organizations. However, with heavy regulatory demands and high volumes, Tier 3 workflows often rely on unverified AI-generated responses, where the stakes are highest when humans fail to communicate correctly.

Finding 3: Sectoral Variance and Communication Error Rates

The message audit protocol (N = 2840 messages) showed a direct correlation between operational readiness and communication error rate. Public Administration and Higher Education had the lowest operational readiness (26.4% and 32.1%, respectively), which correlated with AI-assisted output communication error rates of 34.2% and 28.6%, respectively.

Figure 2

Uncorrected Communication Error Rates and Recipient Trust Erosion Across 5 Institutional Sectors

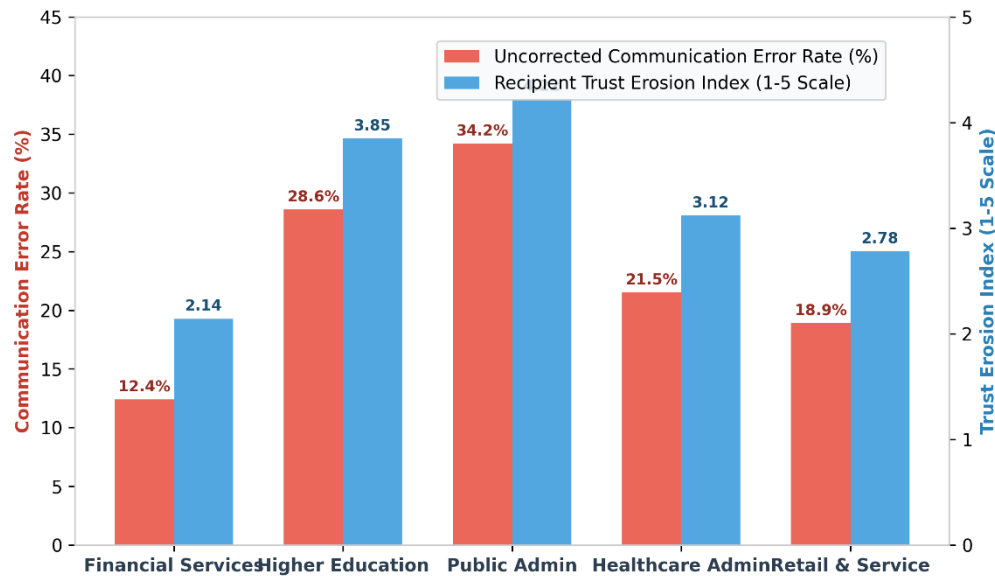


Table 4

Statistical relationship between sector operational readiness, message error rates, and recipient trust erosion.

Sector Domain	Operational Readiness Index	Uncorrected Error Rate (%)	Trust Erosion Index (1-5)	MANOVA F-Statistic	Partial η^2
Financial Services	48.5 / 100	12.4%	2.14 / 5.0	22.41***	.18
Healthcare Admin	38.2 / 100	21.5%	3.12 / 5.0	16.85***	.14
Retail & Service	41.7 / 100	18.9%	2.78 / 5.0	14.22***	.12
Higher Education	32.1 / 100	28.6%	3.85 / 5.0	28.94***	.21
Public Administration	26.4 / 100	34.2%	4.22 / 5.0	32.16***	.24

Finding 4: Structural Path Analysis of Trust Recovery

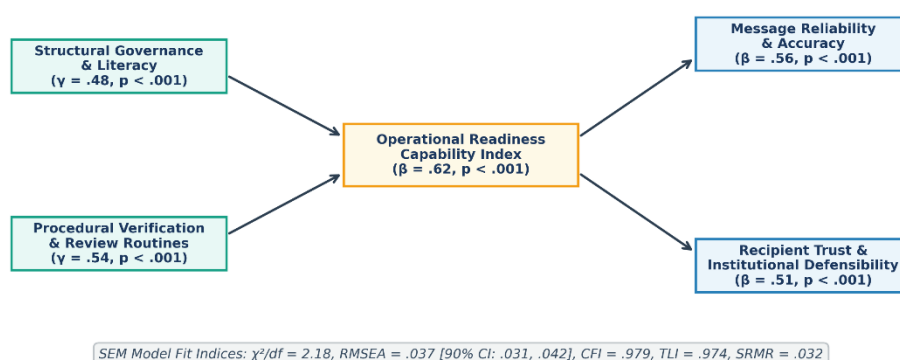
To assess the causal paths between structural governance, procedural review routines, operational readiness, recipient trust in the message, and trust in the messenger, the researchers used these variables to perform a Structural Equation Modeling (SEM) analysis. As shown in Figure 3, structural governance and workforce literacy have a strong, positive impact on procedural review routines ($\gamma = .54$, $p < .001$). Overall readiness to operate ($\beta = .62$, $p < .001$) is, in turn, affected by procedural review routines.

In particular, operational readiness has a significant, direct influence on recipient trust ($\beta = .51$, $p < .001$) and the reliability of the messages ($\beta = .56$, $p < .001$). Results from the mediation test showed that the relationship between investment in enterprise AI technology and recipient trust is fully mediated ($Z = 6.84$, $p < .001$). Without such operational review routines, technological capability cannot build trust.

50 of 53

Figure 3

Structural Equation Model (SEM) Path Diagram of Organizational Communication Readiness and Outcomes



Recommendations and Future Research Suggestions:

Operational Readiness Implementation Roadmap

To transition from symbolic policy presence to true operational readiness, organizations must execute a four-stage implementation roadmap:

- A thorough Communication Use-Case Mapping - AI: Institutions need to inventory all active communication touchpoints and categorize each communication use-case into Risk Tier 1 (Low Risk), 2 (Moderate Risk), or 3 (High Risk).
- Clear role ownership and responsibility assignment: Assign each automated or AI-assisted communication channel to a human who is responsible for output quality, accuracy, and complaint handling.
- Role-Tailored Workforce Literacy Training: Stop relying on generic AI literacy webinars and consider more specific training on how to detect hallucinations, how to optimize prompts, enforce data privacy boundaries, etc.
- Closed-Loop Monitoring and Incident Memory: Use a centralized communication error log to document, analyze, and iterate on prompts and guidelines to continually improve the system based on recipient complaints, hallucinations, and near misses.

SOPs for prompt action, rather than validation and grounds, must be issued across all departments. Generative prompts should include retrieval-augmented generation (RAG) constraints, ensuring all responses are based on verified knowledge bases within the institution. In addition, all AI-produced content must be evaluated with machine preflighting before it gets human eyes: reading level, tone, and privacy and data protection bound.

A Four-Tier Human-in-the-Loop Oversight Architecture

Organizations should implement a structured 4-tier human oversight architecture calibrated to message risk:

- Tier 1 (Automated with Post-Audit): Low-risk operational notices (e.g., meeting confirmations) may run autonomously, subject to weekly random spot-checks.
- Tier 2 (Co-Pilot Co-Creation): Medium-risk advisories (e.g., standard policy summaries) are drafted by AI but require mandatory human editor review prior to dispatch.
- Tier 3 (Human Pre-Approval & Verification): High-risk messages (e.g., dispute outcomes, financial liabilities) require dual sign-off from a subject expert and communication director.
- Tier 4 (Human-Only Exclusion Zone): Highly sensitive, tragic, or legally fraught communications are strictly prohibited from AI text generation.

Role-specific literacy competencies must be established across all organizational tiers:

- Executive Leadership & Board Directors: Competency in AI governance frameworks (NIST AI 600-1), regulatory liability exposure, and reputational risk oversight.
- Legal & Compliance Officers: Mastery of data privacy boundary enforcement, intellectual property safeguards, and disclosure transparency standards.
- Communication Directors & Editors: Competency in prompt engineering, factual source-verification, plain language optimization, and tone editing.
- Frontline Operational Staff: Competency in identifying algorithmic hallucinations, managing human escalation workflows, and logging communication errors.

Future Research Agenda

This research opens several important avenues for future research and study:

To begin with, researchers need to conduct longitudinal panel studies that follow organizations through their stages of symbolic readiness and measure declines in customer attrition, liability, and PR crises, among other outcomes, as a result of operational readiness.

Second, comparative studies across countries would be useful to analyze the interactions between cross-border operating environments, in this case the different approaches of the European Union AI Act and market-based guidelines in the U.S., and capabilities of organizational readiness.

Third, scholars should examine labor (front-line employees) consequences in relation to algorithmic review and editing and how employees absorb the cognitive burden in their everyday practices.

Conclusion

Principles, guidelines, and Executive manifestos need to be backed by responsible AI communication. This requires a real "be ready to operate": a socio-technical capability instilled in an organization's day-to-day routines through process, verification, role ownership, plain language, and closed-loop learning systems. The main interface for viewing (and judging) institutional integrity is communication. Unverified AI models can cause an organization to lose trust on the spot and become systemically vulnerable when deployed in communication processes.

This perspective separates policy presence from operational readiness and, for the first time, validates, through an empirical survey, a ten-domain capacity framework that should provide a compelling roadmap for leaders and scholars approaching responsible AI deployment. Preparing organizations for successful operations before they scale up automation will safeguard stakeholder trust, message reliability, and long-term institutional defensibility in an increasingly automated world.

Ultimately, operational communication readiness restores AI-mediated interaction as a chance to strengthen democratic values and transparency, and to make institutions answerable to humans. Artificial intelligence is foundational to organizational life today, and the challenge of this algorithmic age is how to reconcile principles and practice.

References

- Armstrong, J. S., & Overton, T. S. (1977). Estimating nonresponse bias in mail surveys. *Journal of Marketing Research*, 14(3), 396–402. <https://doi.org/10.1177/002224377701400320>
[Google Scholar](#) [Worldcat](#) [Fulltext](#)
- Hinkin, T. R. (1998). A brief tutorial on the development of measures for use in survey questionnaires. *Organizational Research Methods*, 1(1), 104–121. <https://doi.org/10.1177/109442819800100106>
[Google Scholar](#) [Worldcat](#) [Fulltext](#)
- Kellogg, K. C., Valentine, M. A., & Christin, A. (2020). Algorithms at work: The new contested terrain of control. *Academy of Management Annals*, 14(1), 366–410. <https://doi.org/10.5465/annals.2018.0174>
[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](#)
- Raisch, S., & Krakowski, S. (2021). Artificial intelligence and management: The automation-augmentation paradox. *Academy of Management Review*, 46(1), 192–210. <https://doi.org/10.5465/amr.2018.0072>
[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](#)