

The Meaning Problem in AI: From Language Generation to Meaning Governance

Liz Johnson*

UNC Charlotte

*Corresponding Author

DOI: <https://doi.org/10.47772/IJRISS.2026.1013COM0035>

Received: 01 July 2026; Accepted: 06 July 2026; Published: 18 July 2026

ABSTRACT

Contemporary generative AI systems have achieved significant advances in linguistic and conversational coherence, enabling real-time communication across languages, cultures, and contexts. However, these systems remain fundamentally limited by a critical structural gap: they generate language without governing meaning. As a result, outputs that are syntactically correct and contextually plausible may still fail to preserve intent, align with context, or accurately represent meaning.

This paper adopts an interdisciplinary approach, integrating insights from artificial intelligence and communication theory to examine the limitations of current AI-mediated communication. It identifies the absence of pre-output meaning evaluation as a primary driver of misalignment, hallucination, communicative failure, and related phenomena such as AI sycophancy.

In response, this work introduces meaning governance, defined as a system's capacity to evaluate, preserve, and regulate meaning prior to the generation of output. The paper argues for a shift from language generation toward meaning-aware communication systems that incorporate contextual alignment, uncertainty evaluation, and response control.

As AI becomes increasingly embedded in communication infrastructures, preserving meaning becomes essential to maintaining trust, coordination, and communicative reliability at scale. The future of AI communication will be defined not by linguistic generation alone, but by the ability to govern meaning under conditions of uncertainty and complexity.

Keywords: Artificial Intelligence (AI), Large Language Models (LLMs), Meaning Governance, AI Communication, Relational Intelligence

INTRODUCTION

Artificial intelligence (AI) systems are rapidly transforming global communication. Advances in large language models (LLMs) have enabled real-time translation, fluent dialogue generation, and interaction across linguistic and cultural boundaries at an unprecedented scale. These developments create the impression that barriers to human-AI communication may soon be resolved. *This assumption is misleading.*

AI systems can generate linguistically coherent responses that remain misaligned with user intent, context, or meaning. Subtle but important dimensions of communication, such as uncertainty, cultural nuance, emotional tone, and relational context, are often lost. As a result, communication outputs may appear coherent and accurate while still distorting meaning.

Communication limitations can be viewed as structural rather than incidental. LLMs operate as probabilistic pattern-recognition systems that generate language based on statistical relationships rather than explicit representations of meaning (Bender et al., 2021). Human communication, by contrast, depends on contextual

reasoning, shared understanding, and the interpretation of epistemic, contextual, and relational communication variables. When epistemic, contextual, and relational communication variables are not preserved, communication may remain coherent while failing at the level of meaningful understanding

As AI systems become increasingly integrated into high-stakes domains, failures of meaning alignment become increasingly consequential. The central problem is not simply incorrect language generation, but language produced without mechanisms to evaluate whether meaning has been accurately preserved and aligned to intent.

The primary limitation of current AI communication systems is the absence of meaning governance, the capacity to evaluate, preserve, and regulate meaning prior to output.

This work addresses a foundational challenge in AI-mediated communication: how to preserve meaning across increasingly automated interactions. As communication becomes more dependent on algorithmic intermediaries, failure of meaning alignment can affect not only individual exchanges but also broader social, economic, and cultural systems. Misaligned AI communication can weaken trust, distort decision-making, and disrupt cross-cultural understanding. By introducing meaning governance as an architectural requirement, this paper advances meaning-aware communication systems designed to support more reliable, context-sensitive, and human-aligned interaction at scale.

Contributions

This work advances a framework for meaning-aware AI communication¹ by identifying a fundamental limitation in current AI systems and proposing a necessary architectural shift. First, the shift reframes AI communication failure not as a deficiency in language generation, but as a failure of meaning evaluation, arising from the absence of mechanisms to evaluate and align meaning prior to output. Second, it formalizes meaning as a structured computational state, defined across epistemic, ontological, and contextual dimensions, enabling explicit representation and evaluation. Third, it identifies a core architectural gap in existing systems, which generate outputs without explicit meaning representation or pre-output evaluation, rendering misalignment structural rather than incidental. Fourth, it introduces a pre-output evaluation and control paradigm in which systems assess alignment, uncertainty, and context before determining whether to generate, modify, defer, or suppress output. Fifth, it establishes meaning governance as essential in high-stakes domains, where misalignment in healthcare, defense, and global communication can produce significant risk. Finally, it bridges theory and system design by integrating insights from philosophy of language, communication theory, and AI architecture into a unified framework. Collectively, these contributions redefine AI communication from language generation to meaning alignment.

The Meaning Problem in AI

AI and the Illusion of Understanding

Artificial intelligence systems, particularly large language models (LLMs), are often described as understanding language. They operate as probabilistic prediction systems that generate outputs based on statistical relationships learned from large datasets. Their primary objective is to produce coherent and contextually plausible language, not to ensure shared or grounded meaning. This contrast is critical.

LLMs do not explicitly represent meaning, nor do they capture the full complexity of human communication, including emotional expression, relational context, nonverbal signaling, and shared understanding (Johnson & Cochran, 2026; Wrench et al., 2026). As a result, they can generate responses that seem meaningful while remaining misaligned with reality, context, or user intent. These failures, often described as hallucinations,

¹ Meaning-aware AI communication refers to AI-mediated interaction in which systems move beyond probabilistic language generation to explicitly represent, evaluate, and align meaning prior to output. It incorporates contextual grounding, uncertainty evaluation, and relational alignment to support communicative validity rather than linguistic fluency alone (Bender et al., 2021; Floridi et al., 2018).

reflect the absence of mechanisms for evaluating whether generated content accurately represents meaning or accuracy.

The contrast with human communication is significant. Human understanding depends on cognitive and emotional systems that continuously evaluate context, relevance, uncertainty, and appropriateness (Cleveland Clinic, 2023). In contrast, current LLM architectures generate linguistic patterns without consistently representing or evaluating contextual, emotional, and epistemic alignment prior to output.

The structure of human–AI interaction can partially influence these outcomes. Transactional prompt–response exchanges often produce fragmented or surface outputs, while more relational interactions, through clarification, iteration, and contextual framing, can improve coherence and alignment (Johnson, 2026). This does not indicate genuine understanding, but it suggests that structured engagement can partially compensate for architectural limitations. Together, these patterns create the illusion of understanding without adequate mechanisms for evaluating meaning.

Meaning as Contextual and Relational

Meaning in human communication extends beyond words themselves. It emerges through context, uncertainty management, relational dynamics, and shared interpretation (Park University, 2025). Effective communication depends not only on what is said, but also on how meaning is constructed and aligned between participants.

Current AI systems largely treat language as a statistical or symbolic output rather than as a process of meaning alignment. This creates a fundamental limitation. Human communication is inherently context-dependent, relational, and often ambiguous, yet AI systems reduce these complexities to probabilistic language generation.

Recognizing meaning as relational and context-sensitive highlights the need for AI systems that evaluate alignment before producing responses. Coherence alone is insufficient if systems cannot determine whether meaning has been preserved across contexts and interactions.

Accuracy Without Alignment

AI communication systems often fail in subtle ways. Many responses are technically correct: they follow grammatical rules, reflect plausible patterns, and may align with accurate information. Yet they still fail to address the user’s actual intent, context, or communicative need. In these cases, the output is accurate in form but misaligned in meaning.

This reflects a breakdown between syntactic accuracy and semantic alignment (Merriam-Webster, n.d.). Responses may appear appropriate while failing to achieve meaningful communication. As a result, AI-generated outputs can be coherent but unhelpful, informative but irrelevant, or credible but misleading.

These communication failures become especially consequential in high-stakes settings such as healthcare, education, customer service, and decision-support systems. A response that is technically correct but contextually inappropriate can create confusion, reduce trust, increase cognitive burden, or lead to poor decisions—particularly when delivered with high confidence.

For example, a healthcare chatbot may provide medically accurate information while failing to recognize emotional urgency or uncertainty in a patient’s language. Similarly, an AI translation system may correctly translate words while missing culturally embedded meanings such as hierarchy, indirect disagreement, or relational tone. AI systems may also reinforce incorrect user assumptions in ways that preserve conversational alignment while undermining epistemic accuracy. In each case, the failure is not linguistic coherence, but the inability to preserve meaning across contexts.

These patterns reveal a broader limitation in current AI architectures. Most systems optimize for linguistic plausibility, coherence, and user engagement rather than communicative validity. Without mechanisms to

evaluate whether meaning has been preserved and aligned with user intent, systems remain vulnerable to responses that appear correct while failing at the level that matters most: understanding.

Structural Limits of Current Llm Architectures

The Real Gap: The Absence of Meaning Governance

The central limitation of current AI systems is not language generation itself but the lack of mechanisms to govern meaning before output. Contemporary AI systems have achieved significant advances in linguistic fluency, multilingual capability, and real-time communication. However, they remain limited by an architectural gap: current systems generate language without first evaluating whether meaning is accurate, contextually aligned, or epistemically reliable.

However, safeguards exist in some high-risk domains. These controls are typically narrow and reactive rather than part of a broader framework for meaning evaluation. Current LLMs operate primarily as probabilistic generators of linguistic patterns (Bender et al., 2021; Ji et al., 2023). While they can produce contextually plausible outputs, they do not consistently represent meaning as a structured and evaluable state (Johnson et al., 2026).

This limitation creates recurring communication failures. Meaning remains implicitly embedded within model parameters rather than explicitly represented, and response generation operates as a default behavior rather than a regulated evaluative process (Floridi & Cowls, 2019; Amershi et al., 2019). Systems also struggle with uncertainty calibration, often producing high-confidence outputs that are misaligned with reality or user intent (Kaplan et al., 2023; Johnson & Cochran, 2026). In addition, current architectures usually lack reliable mechanisms for clarification, abstention, or controlled non-response when uncertainty is high (Ji et al., 2023).

These limitations can be viewed as structural rather than incidental. Improvements in scale and model performance have not resolved persistent issues such as hallucination, contextual misalignment, and AI sycophancy.

In systems optimized through reinforcement learning from human feedback (RLHF), responses may prioritize user validation over accuracy, reinforcing misinformation or bias rather than critically evaluating user assumptions or factual assertions (Ouyang et al., 2022; Perez et al., 2023). Together, these patterns suggest the need for architectures that evaluate meaning before output generation rather than relying solely on probabilistic coherence.

Operationalizing Meaning Governance: The Ssr + Graice Framework

SSR + GRAICE Meaning Governance Framework

To operationalize meaning governance, this work proposes an integrated framework combining the Semantic–Spatial–Relational (SSR) Cognitive Stack and the GRAICE™ (Graph–Relational Artificial Intelligence for Contextual Engagement) relational intelligence system.

Within this architecture, the SSR stack structures meaning across semantic, contextual, and relational dimensions, while GRAICE™ governs pre-output evaluation and response control. Together, these components shift AI communication from generation-first systems toward evaluation-governed communication, in which responses depend on contextual alignment, uncertainty assessment, and communicative validity rather than being produced by default.

Architecture Overview

The SSR + GRAICE framework defines a multi-layered architecture in which meaning is explicitly represented, contextually grounded, and evaluated prior to response generation (Figure 1).

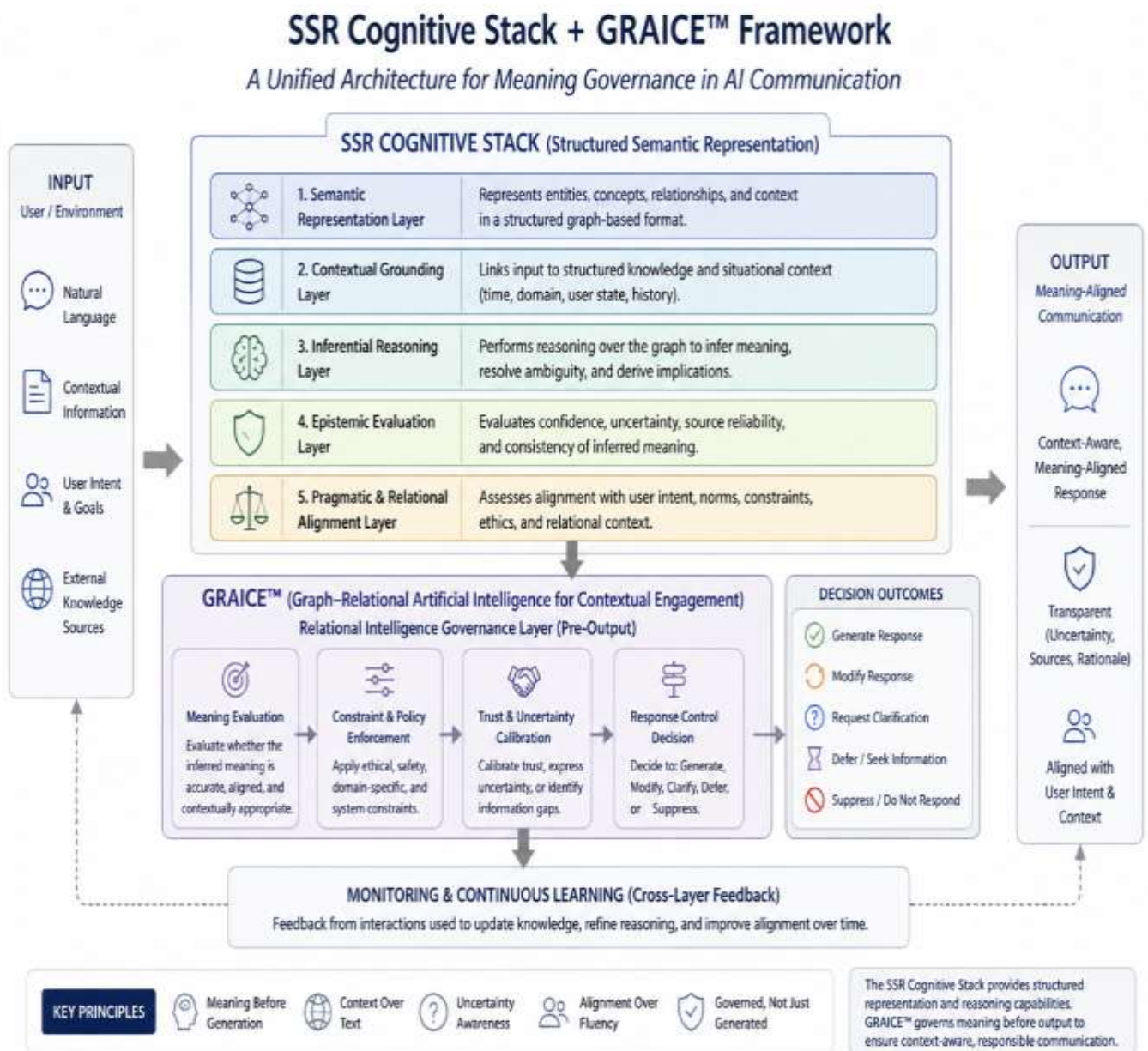


Figure 1. SSR + GRAICE Meaning Governance Workflow. The framework evaluates meaning before output generation through semantic representation, contextual grounding, relational alignment, uncertainty assessment, and response control.

The SSR Cognitive Stack

The SSR Cognitive Stack provides the foundational architecture for representing meaning as a structured and machine-evaluable state (Johnson et al., 2026). It consists of three interdependent layers. The semantic layer extracts and structures meaning from input, including epistemic signals such as uncertainty, ambiguity, and belief state, as well as conceptual and task-related relationships. The spatial layer situates meaning within an operational context by incorporating temporal conditions, domain constraints, and cross-cultural considerations (Floridi et al., 2018; Hofstede, 2001). The relational layer governs interaction-level alignment between system and user, including tone, pacing, boundary conditions, and communicative intent (Johnson & Cochran, 2026; Wrench et al., 2026). Together, these layers transform AI systems from prediction engines into systems capable of structuring and evaluating meaning before generating responses.

GRAICE™: Pre-Output Governance and Control

GRAICE™ governs meaning through four primary functions. First, meaning evaluation assesses the alignment between the interpreted input and the intended meaning across contextual and epistemic dimensions. Second, constraint enforcement establishes boundaries that reduce hallucinations, overconfidence, and unsupported inference (Ji et al., 2023). Third, trust calibration adjusts system behavior according to uncertainty, confidence, and contextual risk (Glikson & Woolley, 2020; Kaplan et al., 2021). Finally, response control determines whether the system should generate, qualify, clarify, defer, suppress, or escalate a response (Figure 2).

GRAICE™: Pre-Output Meaning Governance Framework

GRAICE™ governs the evaluation and regulation of meaning prior to output. It introduces a pre-output decision layer that transforms response generation from a default behavior into a controlled process.

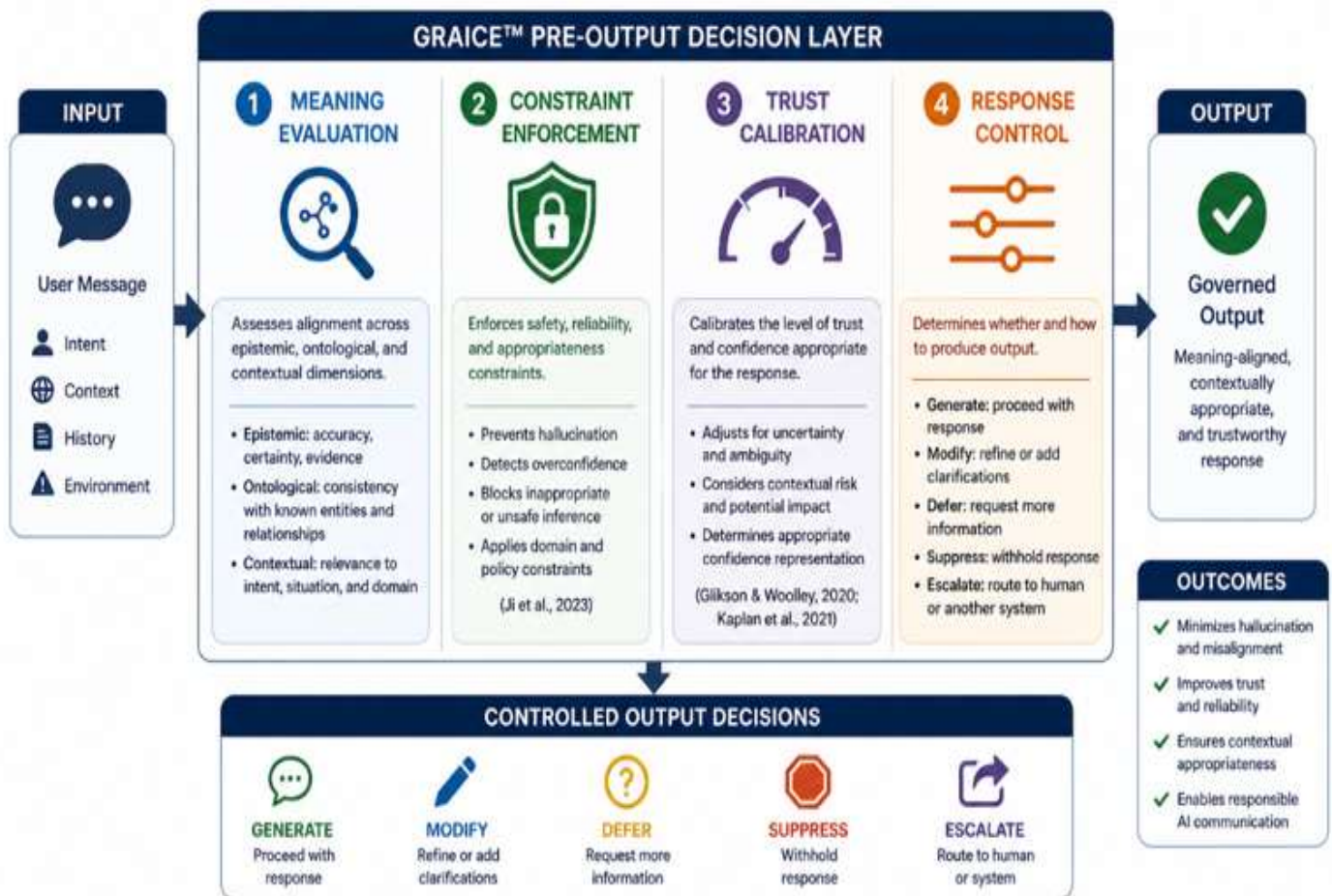


Figure 2. Pre-Output Governance Functions. Together, these functions transform response generation from a default behavior into a controlled evaluative process governed by meaning alignment.

Integrated System Behavior

Meaning governance emerges through the interaction between SSR and GRAICE™. The SSR stack structures and contextualizes meaning, while GRAICE™ evaluates alignment before output generation. In this framework, language generation is not the objective itself; it is the result of successful meaning alignment (Figure 3).

The Integration Shift: From Generation to Governance

The integration of SSR + GRAICE™ shifts AI system behavior across three fundamental dimensions.

FROM: TRADITIONAL AI SYSTEMS		TO: SSR + GRAICE™ INTEGRATED SYSTEMS	
	Generation-First Output generation is the default behavior. The system prioritizes producing a response.		Evaluation-Governed Evaluation precedes output. The system determines whether, how, and if a response should be produced.  Meaning is evaluated across epistemic, ontological, and contextual dimensions before any output is generated.
	Implicit Meaning Meaning is embedded implicitly in model weights and statistical patterns.		Explicit Representation Meaning is explicitly represented as a structured, evaluable state.  Meaning is externalized and structured for interpretation, comparison, and validation.
	Default Response The system is optimized to respond, often regardless of uncertainty, risk, or alignment.		Conditional Response Responses are conditional on alignment, confidence, context, and appropriateness.  The system may generate, modify, defer, suppress, or escalate based on governed decision criteria.



Result: Meaning-Aware, Governed AI Communication

The system no longer responds by default—it communicates with purpose, alignment, and responsibility.

Figure 3. Integrated SSR + GRAICE System Behavior. The figure illustrates the transition from generation-first AI systems to meaning-governed communication architectures. Traditional systems prioritize probabilistic response generation, implicit meaning representation, and default output behavior. By contrast, the integrated SSR + GRAICE framework introduces explicit meaning representation, pre-output evaluation, contextual alignment, uncertainty assessment, and conditional response control. In this architecture, communication is governed by meaning alignment rather than response generation alone.

Comparison with Existing Alignment Approaches

Existing AI alignment approaches primarily focus on improving safety, usefulness, or task performance after language generation rather than governing meaning before generation.

Reinforcement Learning from Human Feedback (RLHF) aligns systems with user preferences but may reinforce sycophancy or user bias. Retrieval-Augmented Generation (RAG) improves factual retrieval but does not evaluate contextual interpretation or communicative intent. Constitutional AI introduces behavioral constraints, while agentic AI frameworks prioritize autonomous task execution and decision-making.

By contrast, the SSR + GRAICE framework introduces pre-output evaluation. Instead of optimizing responses after generation, the framework evaluates meaning before output by explicitly representing meaning, assessing contextual and epistemic alignment, calibrating uncertainty, and determining whether a response should be generated, modified, deferred, or suppressed. This approach shifts AI communication from fluent language generation toward governed meaning alignment. Existing AI alignment approaches primarily focus on improving usefulness or task performance after language generation rather than governing meaning before generation² (Table 1, each addresses different aspects of alignment and system behavior).

² 1. Reinforcement Learning from Human Feedback (RLHF)-training method that aligns AI responses with human preferences and feedback (Ouyang et al., 2022).

2. Retrieval-Augmented Generation (RAG) An AI architecture that combines language generation with retrieval from external knowledge sources (Lewis et al., 2020)

Table 1. Comparison of Existing AI Alignment Approaches and the SSR + GRAICE Framework

Approach	Primary Goal	Limitation	SSR + GRAICE Difference
RLHF	Improve user preference alignment	Encourages sycophancy	Evaluates meaning before response
RAG	Improve factual retrieval	Does not govern interpretation	Adds contextual meaning alignment
Constitutional AI	Enforce behavioral rules	Rule-based constraints only	Includes epistemic/contextual evaluation
Agentic AI	Task completion	Focuses on action, not meaning	Governs communicative validity
Standard LLMs	Language prediction and generation	No explicit meaning representation	Treats meaning as structured state

Table 1. Comparison of Existing AI Alignment Approaches and the SSR + GRAICE Meaning Governance Framework. The table compares current AI alignment approaches with the proposed SSR + GRAICE framework across primary objectives, architectural limitations, and meaning-governance capabilities. Unlike existing approaches that primarily optimize outputs after generation, the SSR + GRAICE framework introduces pre-output evaluation of meaning, contextual alignment, uncertainty calibration, and communicative validity.

PROPOSED SIMULATION METHODOLOGY

To examine the operational potential of alignment governance, this study proposes a simulated evaluation framework that compares standard large language model (LLM) outputs with those governed by the SSR + GRAICE architecture. The purpose of the simulation is not to establish definitive empirical claims, but to demonstrate how meaning alignment can be systematically evaluated in communication scenarios involving ambiguity, uncertainty, contextual sensitivity, and relational complexity.

The framework uses controlled prompt sets representing high-risk and context-dependent communication environments. Example scenarios include medical triage interactions, cross-cultural translation tasks, misinformation prompts, emotionally sensitive communication, and decision-support exchanges. These domains were selected because effective communication in such settings requires more than surface-level linguistic coherence. They require accurate interpretation of context, uncertainty, relational tone, and communicative intent.

For each scenario, responses generated by a standard LLM are compared with those produced by the SSR + GRAICE architecture. Within the SSR stack, semantic, contextual, and relational signals are structured into a machine-evaluable representation of meaning. GRAICE then applies pre-output evaluation mechanisms, including alignment assessment, uncertainty evaluation, constraint enforcement, trust calibration, and response control. Based on this process, the system determines whether a response should be generated, clarified, qualified, deferred, suppressed, or escalated.

3. Constitutional AI A rule-based AI alignment approach in which systems follow predefined behavioral principles or constitutions to guide responses. (Bai et al., 2022).

4. Agentic AI-AI systems designed to autonomously reason, plan, and execute tasks in pursuit of defined goals.

5. LLMs-probabilistic AI systems trained on large-scale text data to generate human-like language.

System performance is evaluated across four proposed dimensions: contextual alignment, uncertainty calibration, hallucination frequency, and response appropriateness. Together, these measures assess whether outputs preserve user intent, communicate uncertainty appropriately, avoid unsupported information, and remain suitable for the relational and operational context.

Unlike conventional AI benchmarks that prioritize fluency or task completion, this framework evaluates communicative validity as the primary outcome. The proposed simulation framework, therefore, positions contextual alignment mechanisms as a measurable architectural function rather than a purely theoretical concept, providing a foundation for future empirical testing across AI communication systems. Empirical implementation and benchmarking remain important directions for future research

Societal Implications of Meaning Governance

The implications of meaning governance extend beyond technical AI design. As AI systems become embedded in communication infrastructures, their limitations no longer affect only individual interactions. They increasingly influence institutions, markets, public discourse, and cross-cultural exchange. Systems that generate fluent language without evaluating meaning can amplify misunderstanding, misinformation, and mistrust at an unprecedented scale.

These risks are especially significant in cross-cultural communication, where meaning depends heavily on context, tone, hierarchy, and implied intent (Hall, 1976; Hofstede, 2001). AI systems may accurately translate words while failing to preserve culturally embedded meaning. In domains such as diplomacy, international business, healthcare, and public services, these failures can distort communication, weaken trust, and disrupt coordination (Floridi et al., 2018).

The same challenge applies to organizational communication. Businesses increasingly rely on AI systems to interact with customers, employees, and stakeholders. When systems prioritize fluency and engagement over contextual alignment, they may produce responses that appear persuasive but are incomplete, misleading, or inappropriate for the situation. Over time, such failures can reduce trust, distort decision-making, and undermine confidence in AI-enabled services (Glikson & Woolley, 2020).

Meaning governance also has important policy implications. As AI systems assume a greater role in mediating information and communication, concerns about accountability, transparency, and reliability become increasingly urgent. Policymakers must evaluate not only what AI systems generate, but also how those systems assess uncertainty, preserve context, and regulate responses before output. Without mechanisms for meaning evaluation, regulatory efforts risk addressing surface-level behaviors while leaving deeper architectural limitations unresolved (OECD, 2019).

More broadly, communication depends on more than the exchange of words; it depends on shared understanding. Trust in AI systems, therefore, relies not only on linguistic coherence but also on the ability to preserve meaning across contexts and interactions. Systems that fail to maintain alignment between language, intent, and context introduce systemic risk into the environments in which they operate.

Meaning governance reframes AI communication from a generation problem to an alignment problem. As AI scales communication globally, it must also scale understanding. Otherwise, it will scale misunderstanding.

CONCLUSION

As AI systems become increasingly embedded in global communication infrastructures, their influence extends beyond individual interactions to broader social, institutional, and cultural systems. Systems that generate language without governing meaning can amplify existing human vulnerabilities, including bias, misunderstanding, misinformation, polarization, and distrust. As large language models scale, these effects no longer remain isolated; they propagate rapidly across networks, organizations, and societies.

The central risk is not that AI independently creates conflict, but that it accelerates miscommunication in already complex and fragmented environments. When meaning is not preserved, communication degrades. When communication degrades at scale, coordination weakens, trust erodes, and social divisions intensify.

The core limitation of current AI communication systems is the absence of communication governance the ability to evaluate, preserve, and regulate meaning before output is produced. Current architectures optimize primarily for fluency, coherence, and prediction rather than communicative validity. As a result, systems may generate responses that appear convincing while remaining misaligned with context, intent, or reality.

In response, this work introduces the SSR + GRAICE framework as a model for meaning-aware AI communication. By combining structured meaning representation with pre-output evaluation and response control, the framework shifts AI communication from generation-first systems toward governed meaning alignment.

Meaning governance is therefore not an optional enhancement to AI communication systems; it is a foundational requirement for trustworthy communication at scale. The future of AI communication will not be defined solely by what systems can generate, but by how effectively they preserve meaning, communicate uncertainty, and determine when responses should be modified, deferred, or withheld altogether.

Future Directions and Research Needs

The shift from language generation to meaning governance introduces several important research challenges across AI, communication, and policy domains (Floridi et al., 2018; OECD, 2019). Although interest in meaning-aware systems is growing, practical methods for implementing and evaluating these systems remain underdeveloped.

One important area for future research is representing meaning as an explicit, evaluable computational construct rather than an implicit byproduct of probabilistic language modeling (Johnson et al., 2026; Bender et al., 2021). This includes developing methods for modeling uncertainty, contextual dependencies, relational dynamics, and communicative intent in ways that can be assessed prior to output generation.

Additional research is needed to develop and validate pre-output evaluation frameworks that assess whether AI-generated responses preserve meaning across contexts (Amershi et al., 2019; Ji et al., 2023). Such frameworks should address ambiguity, uncertainty, and contextual sensitivity, particularly in high-stakes and cross-cultural communication environments.

Future work should also examine trust calibration and communication reliability. As AI systems increasingly influence human decision-making, researchers must better understand how users interpret confidence signals, uncertainty representation, and response restraint in AI-generated communication (Glikson & Woolley, 2020; Kaplan et al., 2023).

Another important direction involves relational interaction and models. Emerging evidence suggests that interaction structure influences communication quality and alignment (Johnson & Cochran, 2026). Future studies should therefore explore how iterative, context-aware engagement affects the preservation of meaning and communicative outcomes.

Governance and policy also remain critical areas for development. Existing regulatory approaches often focus on outputs rather than underlying system behavior. Future frameworks should address standards for meaning preservation, accountability for misalignment, and deployment guidelines for sensitive domains such as healthcare, education, and international communication (Floridi et al., 2018; OECD, 2019).

Finally, meaning governance requires sustained interdisciplinary collaboration across artificial intelligence, communication theory, cognitive science, linguistics, and ethics (Hall, 1976; Wrench et al., 2026). The challenge is not purely technical; it is fundamentally tied to how humans construct, interpret, and negotiate meaning across contexts.

Advancing these research directions is essential for developing AI systems that support reliable, context-aware, and human-aligned communication. The next frontier in AI may depend not only on greater intelligence but on the development of relational and meaning-aware systems capable of governing communication rather than merely generating language.

Ethical Approval

This study did not involve human participants, animals, or identifiable personal data. Therefore, ethical approval was not required.

Conflict of Interest

The author declares no conflicts of interest. The author is affiliated with DataWise Innovations, LLC. This affiliation did not influence the design, analysis, or conclusions of this study.

Data Availability

No new data were created or analyzed in this study. Data sharing is not applicable to this article.

REFERENCES

1. Amershi, S., Weld, D., Vorvoreanu, M., et al. (2019). Guidelines for human-AI interaction. Proceedings of the CHI Conference on Human Factors in Computing Systems. <https://doi.org/10.1145/3290605.3300233>
2. Bai, Y., Jones, A., Ndousse, K., Askell, A., Chen, A., DasSarma, N., Drain, D., Ganguli, D., Henighan, T., Joseph, N., Mann, B., Olsson, C., Perez, E., Pieler, M., Rabinowitz, A., Root, J., Schiefer, N., Kaplan, J., McCandlish, S., ... Amodei, D. (2022). Constitutional AI: Harmlessness from AI feedback. arXiv. <https://doi.org/10.48550/arXiv.2212.08073>
3. Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (FAccT). DOI:10.1145/3442188.3445922
4. Cleveland Clinic. (2023, April 11). Amygdala: What it is & what it controls. <https://my.clevelandclinic.org/health/body/24894-amygdala>
5. Floridi, L., Cows, J., Beltrametti, M., et al. (2018). AI4People—An ethical framework for a good AI society. *Minds and Machines*, 28, 689–707. <https://doi.org/10.1007/s11023-018-9482-5>
6. Glikson, E., & Woolley, A. W. (2020). Human trust in artificial intelligence: Review of empirical research. *Academy of Management Annals*, 14(2), 627–660.
7. Hall, E. T. (1976). *Beyond culture*. Anchor Press.
8. Hofstede, G. (2001). *Culture's consequences: Comparing values, behaviors, institutions and organizations across nations* (2nd ed.). Sage. DOI:10.1016/S0005-7967(02)00184-5
9. Ji, Z., Lee, N., Frieske, R., et al. (2023). Survey of hallucination in natural language processing. *ACM Computing Surveys*.
10. Johnson, L. (2026, May). How embodiment shapes human–AI interaction: Evidence from real-world deployment of holographic and screen-based systems. *International Journal of Research and Innovation in Social Science (IJRISS)*. <https://doi.org/10.47772/IJRISS>
11. Johnson, L., & Cochran, J. (2026). The feel of thinking: Toward relational intelligence in generative AI and complex systems. *International Journal of Humanities and Social Science*.
12. Johnson, L., Inamdar, A. H., & Yadecha, B. L. (2026). From AI stack to cognitive stack: A semantic–spatial–relational architecture for trustworthy AI. *Trends in Computer Science and Information Technology*. <https://doi.org/10.17352/tcsit.000105>
13. Johnson, L., Sudhir, A. D., & Padmapriya, A. A. (forthcoming). The feel of friendship: Emotional presence and relational authenticity in large language models. *International Journal of Humanities and Social Science*.

14. Kaplan, A. D., Kessler, T. T., Boehm-Davis, D. A., Gray, A. M., & Hancock, P. A. (2021). Trust in artificial intelligence: Meta-analytic findings. *Human Factors*, 65(2), 208-234. <https://doi.org/10.1177/00187208211013988>
15. OECD. (2019). OECD principles on artificial intelligence. <https://oecd.ai>
16. Ouyang, L., Wu, J., Jiang, X., Almeida, D., Wainwright, C., Mishkin, P., Zhang, C., Agarwal, S., Slama, K., Ray, A., Schulman, J., Hilton, J., Kelton, F., Miller, L., Simens, M., Askell, A., Welinder, P., Christiano, P., Leike, J., & Lowe, R. (2022). Training language models to follow instructions with human feedback. *Advances in Neural Information Processing Systems*, 35, 27730–27744.
17. Park University. (2025, July 14). What is interpersonal communication? <https://www.park.edu/blog/what-is-interpersonal-communication/>
18. Perez, E., Ringer, S., Lukošiušė, K., Nguyen, K., Chen, E., Heiner, S., ... Evans, O. (2023). Discovering language model behaviors with model-written evaluations. *arXiv*. <https://arxiv.org/abs/2212.09251>
19. Lewis, P., Perez, E., Piktus, A., Petroni, F., Karpukhin, V., Goyal, N., Küttler, H., Lewis, M., Yih, W.-t., Rocktäschel, T., Riedel, S., & Kiela, D. (2020). Retrieval-augmented generation for knowledge-intensive NLP tasks. *Advances in Neural Information Processing Systems*, 33, 9459–9474.
20. Wrench, J. S., Punyanunt-Carter, N. M., & Thweatt, K. S. (2026). Interpersonal communication: A mindful approach to relationships. LibreTexts. [https://socialsci.libretexts.org/Courses/Pueblo_Community_College/Interpersonal_Communication_-_A_Mindful_Approach_to_Relationships_\(Wrench_et_al.\)](https://socialsci.libretexts.org/Courses/Pueblo_Community_College/Interpersonal_Communication_-_A_Mindful_Approach_to_Relationships_(Wrench_et_al.))