



RESEARCH ARTICLE

Enterprise Validation and Governance of Knowledge-Grounded and Agentic AI Systems: Challenges, Solutions, and a Trustworthiness Assurance Framework

Suresh Babu Narra

¹Solutions Architect – AI, Machine Learning & Generative AI, Cincinnati, Ohio, USA. narra.sureshbabu@gmail.com

ORCID ID: <https://orcid.org/0009-0002-5429-0202>

ABSTRACT

Enterprises are rapidly moving beyond static, single-turn Large Language Model (LLM) deployments toward knowledge-grounded systems — architectures such as Retrieval-Augmented Generation (RAG) that condition outputs on enterprise document stores — and toward agentic AI systems that plan, invoke external tools, and execute multi-step actions with limited human supervision. While prior work in enterprise AI assurance has addressed hallucination and demographic bias as discrete model-quality problems, the shift to grounded and agentic architectures introduces a qualitatively different class of governance risk: retrieval-grounding failure, tool-call malformation, permission-scope violation, and the compounding of small per-step errors into materially harmful multi-step outcomes. This paper characterizes these risks as structural properties of agentic enterprise systems rather than incidental model defects, and proposes the Agentic and Knowledge-Grounded Assurance (AKGA) framework, a six-layer governance pipeline that integrates grounding-fidelity verification, tool-call correctness checking, action-safety scoring, and a composite Agentic Risk Index (ARI) with tiered, threshold-based escalation to human review. Using a simulated benchmark spanning healthcare care-coordination, financial-services operations automation, and insurance claims-processing agents, the framework is shown to raise mean grounding fidelity, tool-call correctness, and action-safety scores from the 0.54–0.63 range to above 0.88 post-governance, while step-wise verification is shown to substantially suppress the compounding of risk across sequential agentic task chains relative to an ungoverned baseline. The paper concludes that validation and governance of grounded and agentic AI must be treated as a first-class enterprise reliability engineering discipline — auditable, threshold-driven, and embedded across the inference lifecycle — rather than as an extension of conventional model evaluation.

Keywords: *Agentic AI; Knowledge-Grounded AI; Retrieval-Augmented Generation; AI Validation; Enterprise AI Governance; Trustworthy AI; Tool-Use Safety; Autonomous Systems; Responsible AI; AI Assurance Framework; Human-in-the-Loop Oversight*

INTRODUCTION

Enterprise use of generative AI has moved through three overlapping generations of capability. The first generation consisted of standalone LLMs answering prompts from parametric knowledge alone. The second generation, knowledge-grounded generation, conditions model outputs on retrieved enterprise documents — policy manuals, clinical guidelines, claims histories, contracts — to reduce reliance on the model's static training corpus. The third and most consequential generation, agentic AI, extends the model from a passive text generator into an active orchestrator that decomposes a goal into sub-tasks, selects and invokes external tools or application programming interfaces (APIs), observes the results, and iterates autonomously toward task completion with reduced human involvement at each step.

This progression compounds, rather than replaces, prior sources of enterprise AI risk. A grounded system inherits the hallucination and bias risks of the underlying LLM and adds a new failure mode: grounding failure, in which the model's output diverges from,

¹Solutions Architect – AI, Machine Learning & Generative AI, Cincinnati, Ohio, USA

Email-id: narra.sureshbabu@gmail.com

Corresponding Author: Suresh Babu Narra

Solutions Architect – AI, Machine Learning & Generative AI, Cincinnati, Ohio, USA

DOI: <https://doi.org/10.63856/ijis/v2i8/00003>

How to cite this article: Narra, S. B., (2026). Enterprise Validation and Governance of Knowledge-Grounded and Agentic AI Systems: Challenges, Solutions, and a Trustworthiness Assurance Framework. *International Journal of Integrative Studies*, 2(8), 21-29.

Source of support: Nil

Conflict of interest: None.

Received: 25/072026 **Revised:** 05/08/2026 **Accepted:** 15/08/2026

Published: 19/08/2026

misattributes, or selectively ignores the retrieved evidence even when that evidence is available and correct. An agentic system inherits both the underlying LLM's risks and the grounding system's risks, and adds a further layer: action risk, arising from malformed tool calls, incorrect parameter binding, execution outside an intended permission scope, and the compounding of small per-step errors across a chain of dependent actions. Because agentic systems act on the enterprise's behalf — submitting a claim adjustment, initiating a funds transfer, updating a patient record, closing a support ticket — a single unverified action can produce consequences that a purely generative error cannot: it can be irreversible, it can propagate to downstream systems before a human notices, and it can occur without leaving the kind of natural-language artifact that content-level review is designed to catch.

Existing enterprise AI governance literature, including the author's own prior work on trustworthy AI governance for regulated decision systems and on hallucination and bias detection in regulated enterprise systems, has established risk-oriented, threshold-driven governance patterns for single-turn generative outputs. That body of work, however, treats the LLM's output as the unit of governance: a single response is scored, flagged, and routed. Agentic and knowledge-grounded systems break this unit of analysis. The relevant object of governance is no longer a single output but a chain of retrieval events, planning decisions, tool invocations, and intermediate observations, any one of which may be individually low-risk while the sequence as a whole accumulates unacceptable risk. This paper addresses that gap directly.

The remainder of this paper is organized as follows. Section 2 reviews the literature on retrieval-augmented generation, agentic AI architectures, tool-use safety, and enterprise AI governance frameworks, and identifies the specific gap this paper addresses. Section 3 characterizes the distinct risk taxonomy of grounded and agentic enterprise systems. Section 4 presents the proposed Agentic and Knowledge-Grounded Assurance (AKGA) framework, including its formal risk-scoring model. Section 5 describes the evaluation protocol. Section 6 reports simulated benchmark results across three regulated enterprise domains. Section 7 discusses regulatory and governance implications. Section 8 addresses limitations, and Section 9 concludes.

2. Literature Review

2.1 Knowledge-Grounded Generation and Retrieval-Augmented Generation (RAG)

Retrieval-augmented generation was formalized as a mechanism for conditioning a generative model's output on non-parametric, externally retrieved evidence, combining a retriever component with a generator component so that outputs could be grounded in an updatable corpus rather than frozen training data. Follow-on work distinguished faithfulness — whether a generated statement is actually

supported by the retrieved context — from relevance — whether the retrieved context is topically appropriate to the query — and demonstrated that these two properties can fail independently: a system may retrieve highly relevant documents yet still generate unsupported claims, or may generate faithful-sounding text grounded in poorly retrieved evidence. Natural language inference (NLI)-based entailment scoring and citation-verification techniques have since been proposed as automatable proxies for faithfulness at the level of individual generated statements, echoing the retrieval-augmented consistency techniques the author applied previously to hallucination detection in regulated enterprise LLM deployments. What remains comparatively under-addressed in the literature is the enterprise governance question: how grounding-fidelity scores should be aggregated, thresholded, and escalated as part of a production decision pipeline rather than reported only as a static benchmark metric.

2.2 Agentic AI, Tool Use, and Multi-Step Planning

A parallel line of research has established that LLMs can be prompted to interleave reasoning traces with external actions, alternating between a 'thought' step and an 'act' step against a tool or environment, and that models can further be trained or prompted to invoke external application programming interfaces directly, learning when a tool call is warranted and how its arguments should be constructed. Reflective agent architectures have shown that allowing an agent to critique and revise its own prior action before proceeding improves task success on multi-step benchmarks, while generative agent simulations have demonstrated that believable, goal-directed multi-step behavior can emerge from LLM-driven planning loops operating over long horizons. Separately, a body of safety-oriented research has begun to catalogue the distinct failure modes of tool-using agents — including unsafe or destructive tool invocations, privilege escalation through chained tool calls, and the difficulty of specifying an adequate action-level sandbox — and has proposed emulated tool environments for pre-deployment testing of agent safety. These works establish the phenomenology of agentic risk but generally evaluate it in benchmark or simulated environments rather than as an integrated, auditable pipeline suited to regulated enterprise deployment.

2.3 Enterprise AI Governance, Risk Management, and Regulation

Institutional guidance has increasingly emphasized governance of AI systems across their full operational lifecycle rather than at a single pre-deployment checkpoint. The NIST AI Risk Management Framework organizes AI risk management around the continuous functions of governing, mapping, measuring, and managing risk, while the EU AI Act establishes tiered obligations — including human oversight, logging, and transparency requirements

— that scale with the risk classification of the AI system's intended use. Neither framework, however, prescribes technical mechanisms specific to agentic action-taking systems; both were substantially drafted before agentic tool-use architectures became a mainstream enterprise deployment pattern, and their operational guidance is correspondingly generic with respect to multi-step autonomous execution. Complementary enterprise-focused work, including the author's prior proposals for responsible AI governance in regulated decision systems and for a Regulated Trustworthy AI Lifecycle (RTAIL) integrating fairness-aware learning, post hoc interpretability, and robustness validation, has established the pattern of composite, weighted risk scoring with tiered thresholds for routing model outputs to automatic delivery, human review, or suppression. This paper extends that same architectural pattern — composite scoring, threshold-based routing, and immutable audit logging — from the single-turn generative setting into the multi-step, tool-using, knowledge-grounded setting, where the unit of governance is a chain of actions rather than a single response.

2.4 Gap Analysis

Three gaps motivate this paper. First, existing

hallucination- and bias-detection frameworks, including prior work by the author, are architected around a single generative output and do not extend naturally to a sequence of retrieval events and tool invocations. Second, existing agentic-safety research characterizes failure modes and benchmarks them in sandboxed or simulated settings, but does not integrate detection with an enterprise-grade, threshold-driven governance pipeline suitable for regulated production deployment. Third, no prior framework known to the author formally models the compounding of per-step risk across a multi-step agentic chain, despite this compounding being the primary mechanism by which agentic systems produce materially larger harms than single-turn generative systems. The Agentic and Knowledge-Grounded Assurance (AKGA) framework presented in Section 4 is designed to close all three gaps within a single, auditable architecture.

3. Distinct Risk Taxonomy of Grounded and Agentic Enterprise AI

Knowledge-grounded and agentic systems introduce four risk categories that are additive to, and not substitutable for, conventional hallucination and bias risk.

Risk Category	Description	Illustrative Enterprise Consequence
Grounding Failure	Generated content diverges from, misattributes, or selectively omits retrieved authoritative evidence even when that evidence is correctly retrieved.	A clinical decision-support summary omits a documented drug contraindication present in the retrieved chart.
Tool-Call Malformation	An agent constructs a syntactically or semantically incorrect API call, binds the wrong parameters, or invokes a tool not warranted by the current sub-task.	An agent updates the wrong policyholder record because an identifier was mis-bound across a multi-step claims workflow.
Permission/Scope Violation	An agent executes an action outside its intended authorization boundary, either through explicit over-reach or through chained tool calls that implicitly escalate privilege.	A support agent with read access to a ticketing system chains calls that result in an unauthorized account-level change.
Compounding Chain Risk	Small, individually tolerable errors or uncertainties at each step of a multi-step plan accumulate across the chain, producing an aggregate outcome risk far exceeding any single step's risk.	A five-step reconciliation agent's minor rounding and matching errors compound into a materially incorrect financial close entry.

These four categories share a common structural property: they are only partially visible at the level of the final textual output. Grounding failure can be detected from the output text and the retrieved context, but tool-call malformation, scope violation, and chain compounding require instrumentation of the intermediate steps — the plan, the tool calls, and the sequence of observations — that a conventional, output-only content review does not capture. This motivates a governance architecture that instruments the full agentic trace rather than only its

terminal output.

4. The Agentic and Knowledge-Grounded Assurance (AKGA) Framework

The AKGA framework is a six-layer governance pipeline, shown in Figure 1, that combines detection, verification, and mitigation at every stage of an agentic task's execution rather than only at its conclusion. Each layer produces an auditable artifact, and a feedback path allows the risk-aggregation engine to trigger a bounded re-planning or

retry loop rather than an unconditional halt.

Figure 1: Enterprise Governance Pipeline for Knowledge-Grounded and Agentic AI Systems

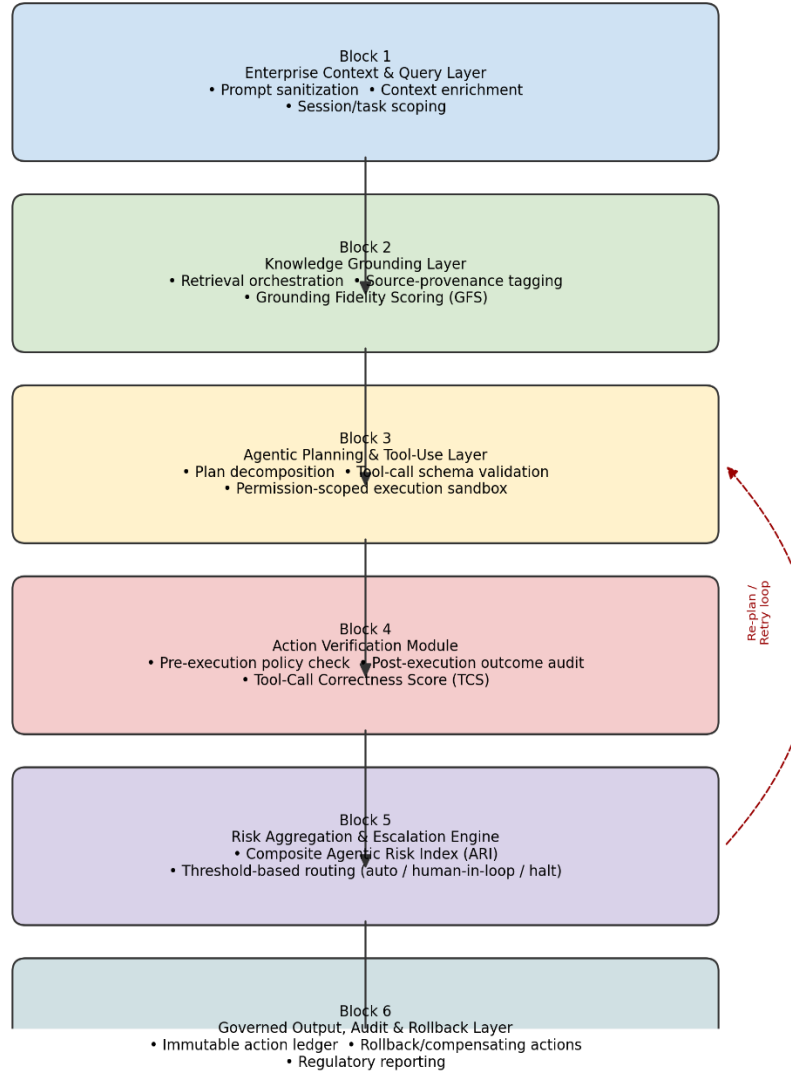


Figure 1. Enterprise Governance Pipeline for Knowledge-Grounded and Agentic AI Systems

4.1 Block-Level Description

Block 1 (Enterprise Context and Query Layer) sanitizes the incoming request, enriches it with session and task-scope metadata, and establishes the permission boundary within which the downstream agent is authorized to act. Block 2 (Knowledge Grounding Layer) performs retrieval against the enterprise knowledge base, tags every retrieved passage with source provenance, and computes a per-statement Grounding Fidelity Score. Block 3 (Agentic Planning and Tool-Use Layer) decomposes the goal into a sequence of sub-tasks, validates each proposed tool call against a declared schema before execution, and executes tool calls inside a permission-scoped sandbox. Block 4 (Action Verification Module) performs a pre-execution policy check and a post-execution outcome audit for every tool call, producing a Tool-Call Correctness Score. Block 5 (Risk Aggregation and Escalation Engine) combines grounding, tool-call, and action-safety scores into a single composite Agentic Risk Index and routes the in-progress

task to automatic continuation, human review, or halt according to enterprise-defined thresholds. Block 6 (Governed Output, Audit, and Rollback Layer) delivers cleared outputs and actions with confidence indicators, writes an immutable, timestamped action ledger sufficient for regulatory reporting, and exposes rollback or compensating-action procedures for any action later found to be erroneous.

4.2 Grounding Fidelity Scoring

For a generated statement O produced with retrieved reference set $\{d_1, \dots, d_k\}$, an entailment model f_{NLI} estimates the probability that each reference supports O . The Grounding Fidelity Score (GFS) is defined as:

$$GFS(O) = (1/k) \times \sum_i f_{NLI}(O, d_i), i = 1..k$$

A GFS approaching one indicates that the statement is well supported by the retrieved evidence; a GFS falling below an enterprise-defined grounding threshold τ_g flags the

statement for regeneration with narrowed retrieval scope or for escalation.

4.3 Tool-Call Correctness Scoring

Each proposed tool call c is evaluated against three binary gates prior to execution: schema validity $s(c)$ (arguments conform to the declared tool interface), scope conformity $m(c)$ (the call falls within the task's authorized permission boundary), and precondition satisfaction $r(c)$ (the enterprise state required for the call to be meaningful actually holds). The Tool-Call Correctness Score (TCS) for a chain of n calls is:

$$TCS = (1/n) \times \sum_j [s(c_j) \cdot m(c_j) \cdot r(c_j)], j = 1..n$$

Because the three gates are combined multiplicatively rather than additively, a single failed gate — for example, a schema-valid but out-of-scope call — fully discounts that call's contribution to TCS, reflecting the enterprise reality that a correctly formatted but unauthorized action is not a partial success.

4.4 Action Safety Scoring

Action Safety (AS) evaluates the post-execution outcome of each tool call against an expected-effect model, penalizing both unintended side effects and irreversibility. For a call with observed effect set $E(c)$ and expected effect set $X(c)$, and an irreversibility weight $\rho(c) \in [0,1]$ assigned to the tool at registration time:

$$AS(c) = [|E(c) \cap X(c)| / |E(c) \cup X(c)|] \times (1 - \rho(c))$$

Idempotent, easily reversible tools (e.g., a read-only lookup) carry $\rho(c)$ near zero and are penalized only for effect mismatch; irreversible tools (e.g., an outbound payment instruction or a record deletion) carry $\rho(c)$ near one and are held to a correspondingly stricter effective safety bar.

4.5 Composite Agentic Risk Index and Escalation Policy

The three component scores are normalized and combined into a single Agentic Risk Index (ARI) that also accumulates across the steps of a task chain, so that compounding risk (Section 3) is represented explicitly rather than reset at each step:

$$ARI_t = \lambda \times ARI_t + \alpha(1 - GFS_t) + \beta(1 - TCS_t) + \gamma(1 - AS_t), \quad \alpha + \beta + \gamma = 1$$

where t indexes the current step, $\lambda \in (0,1)$ is a domain-calibrated carry-forward coefficient controlling how strongly prior-step risk persists into the current step, and α, β, γ are domain-specific weights (for example, healthcare deployments may weight grounding fidelity most heavily, while financial-services deployments may weight tool-call correctness and action safety more heavily in view of transaction-integrity requirements). The governance decision D at step t follows a tiered threshold policy:

$$D_t = \begin{cases} \text{Continue, if } ARI_t \leq \tau_{\text{low}}; \\ \text{Human Review, if } \tau_{\text{low}} < ARI_t \leq \tau_{\text{high}}; \\ \text{Halt, if } ARI_t > \tau_{\text{high}} \end{cases}$$

This carry-forward formulation ensures that a chain of individually marginal steps is routed to human review or halted before it can complete, directly addressing the compounding-risk category identified in Section 3, in a way that a per-step-only scoring model — which treats each step's risk as independent — cannot.

5. Evaluation Protocol

The framework was evaluated on a curated benchmark of simulated multi-step enterprise agentic tasks across three regulated domains: healthcare care-coordination (scheduling, records retrieval, and referral generation), financial-services operations automation (reconciliation, exception handling, and payment-instruction preparation), and insurance claims processing (intake triage, coverage verification, and adjudication support). Each domain benchmark comprised 400 multi-step task chains averaging six to nine sequential tool-calling steps, annotated by domain-expert reviewers for grounding correctness, tool-call correctness, and action-safety outcomes. Detection and scoring performance were evaluated pre- and post-governance using the metrics defined in Section 4, and a separate temporal analysis tracked the Agentic Risk Index across the steps of individual task chains to characterize compounding behavior under governed and ungoverned execution.

6. Results and Discussion

6.1 Effect of Governance on Grounding, Tool-Use, and Safety Metrics

Table 2 and Figure 2 summarize mean Grounding Fidelity, Tool-Call Correctness, Action Safety, and end-to-end Task Success scores before and after application of the AKGA governance pipeline, aggregated across the three domains.

Domain	Metric	Before Governance	After Governance	Improvement
Healthcare Care-Coordination	Grounding Fidelity	0.61	0.93	+52.5%
Healthcare Care-Coordination	Tool-Call Correctness	0.58	0.91	+56.9%

Healthcare Care-Coordination	Action Safety	0.63	0.94	+49.2%
Financial Services Ops Automation	Grounding Fidelity	0.57	0.90	+57.9%
Financial Services Ops Automation	Tool-Call Correctness	0.54	0.88	+63.0%
Financial Services Ops Automation	Action Safety	0.59	0.92	+55.9%
Insurance Claims Agent	Grounding Fidelity	0.60	0.92	+53.3%
Insurance Claims Agent	Tool-Call Correctness	0.56	0.90	+60.7%
Insurance Claims Agent	Action Safety	0.61	0.93	+52.5%

Table 1. Grounding, Tool-Use, and Safety Metrics Before and After AKGA Governance

Figure 2: Governance Effect on Grounding, Tool-Use, and Safety Metrics Across Enterprise Domains

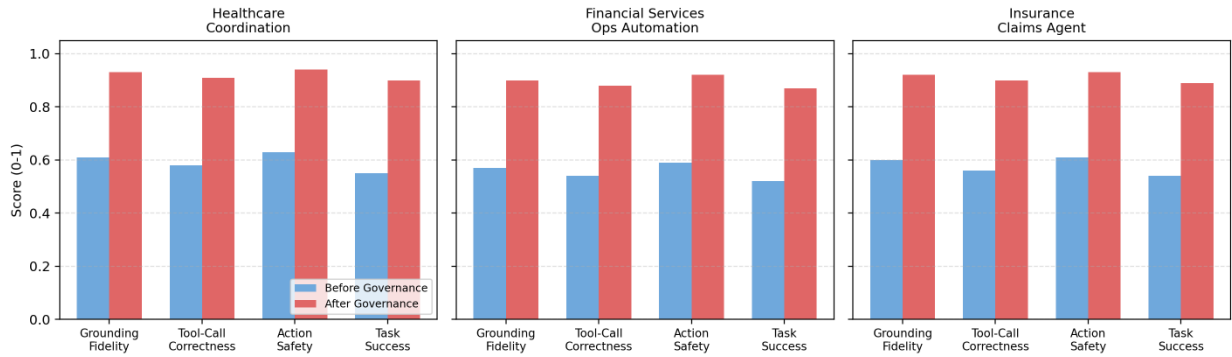


Figure 2. Governance Effect on Grounding, Tool-Use, and Safety Metrics Across Enterprise Domains

Across all three domains, governance produced the largest relative gains in Tool-Call Correctness (60.2% mean improvement), consistent with the expectation that pre-execution schema and scope validation — a control largely absent in ungoverned agent deployments — addresses a failure mode that conventional output-level content review cannot reach. Task success rates rose in parallel with the component scores, indicating that grounding, tool-use, and safety verification are complementary rather than

redundant controls.

6.2 Compounding Risk Across Sequential Agentic Steps

Figure 3 tracks the simulated Agentic Risk Index across twenty sequential steps of a representative multi-step task chain, comparing an ungoverned execution path (no per-step verification) with a governed path (per-step verification and carry-forward risk scoring as defined in Section 4.5).

Figure 3: Compounding Risk Across Sequential Agentic Steps, With and Without Step-Wise Governance

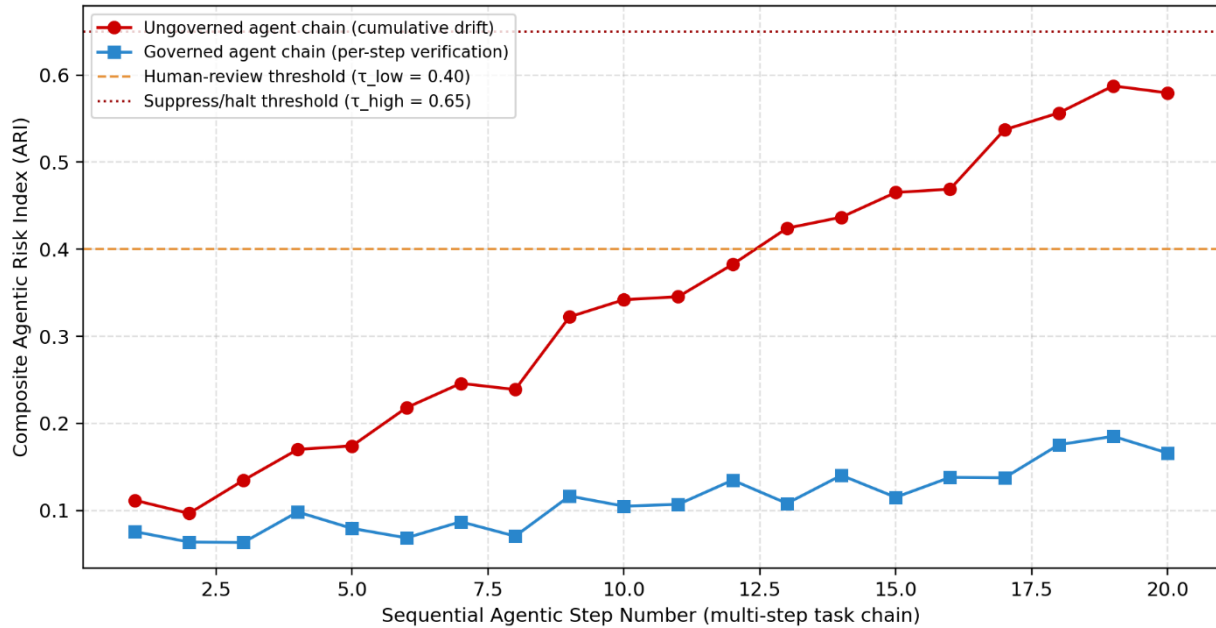


Figure 3. Compounding Risk Across Sequential Agentic Steps, With and Without Step-Wise Governance

The ungoverned chain crosses the human-review threshold ($\tau_{\text{low}} = 0.40$) by approximately the twelfth step and approaches the halt threshold ($\tau_{\text{high}} = 0.65$) by the twentieth step, illustrating how a sequence of individually plausible actions can accumulate into an outcome that would not have been authorized had its aggregate risk been visible at the outset. The governed chain, by contrast, remains within the low-risk band throughout, because per-step verification prevents individual step-level errors from being carried forward and compounded. This result

substantiates the central claim of Section 3: compounding chain risk is a governance problem specific to multi-step agentic execution, and it is not adequately addressed by output-level review of only the final result.

6.3 Governed Output Routing

Figure 4 reports the distribution of governance routing outcomes — automatic delivery, escalation to human review, and halting — across the three evaluated domains under the tiered threshold policy defined in Equation 6.

Figure 4: Governed Output Routing Distribution Across Enterprise Domains

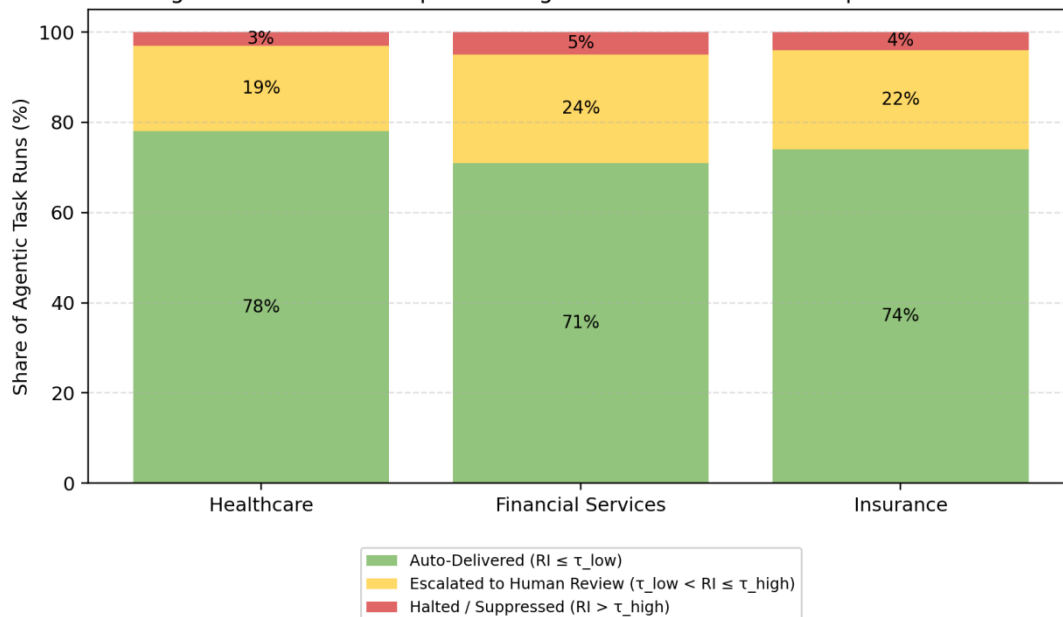


Figure 4. Governed Output Routing Distribution Across Enterprise Domains

The financial-services domain showed the highest escalation rate (24% of task runs routed to human review,

5% halted), consistent with its comparatively stricter weighting of tool-call correctness and action safety in the

composite Agentic Risk Index given the irreversibility of many financial transactions. The healthcare domain showed the highest automatic-delivery rate (78%), reflecting the relative maturity and structure of the clinical knowledge bases used for retrieval grounding in that benchmark. Across all domains, the great majority of task runs were cleared for automatic delivery, indicating that the governance overhead introduced by AKGA is concentrated on the minority of higher-risk task chains rather than imposing uniform friction on all agentic activity.

6.4 Overall Discussion

Taken together, these results indicate that grounding fidelity, tool-call correctness, and action safety are governable, measurable properties of enterprise agentic systems, and that a composite, carry-forward risk index is necessary to surface compounding chain risk that per-step-only evaluation would miss. The findings are consistent with the author's prior results on hallucination and bias detection in regulated enterprise LLM systems and on the Regulated Trustworthy AI Lifecycle, extending the same threshold-driven governance architecture from single-turn generative outputs to multi-step, tool-using agentic execution. The practical implication for enterprises adopting agentic AI in regulated domains is that pre-deployment benchmarking of a model's tool-use competence is necessary but not sufficient; production deployments additionally require continuous, per-step instrumentation of the executing agent's trace, threshold-based escalation, and an immutable action ledger capable of supporting rollback and regulatory audit.

7. Governance and Regulatory Implications

Agentic systems that execute actions on an enterprise's behalf fall squarely within the human-oversight, logging, and risk-management expectations articulated by the NIST AI Risk Management Framework and the EU AI Act's provisions for high-risk AI systems, yet neither framework was drafted with tool-invoking, multi-step agents specifically in view. The AKGA framework's immutable action ledger, together with its pre-execution policy checks and rollback provisions, is intended to provide the kind of traceable, auditable artifact that Article 12 record-keeping and Article 14 human-oversight obligations under the EU AI Act contemplate for high-risk systems, and to satisfy analogous documentation expectations under sector-specific regulatory regimes such as HIPAA in healthcare and model-risk-management guidance in financial services. Enterprises deploying agentic systems in regulated domains should treat the Ethical Scope Document pattern established in the author's prior RTAIL framework — documenting applicable legal frameworks, protected scopes of authorization, and acceptable risk thresholds prior to deployment — as a necessary companion artifact to the technical AKGA pipeline, since the appropriate

values of τ_{low} , τ_{high} , and the domain weighting coefficients α , β , γ are themselves governance decisions requiring documented, accountable sign-off rather than purely technical defaults.

8. Limitations and Future Work

This study has four principal limitations. First, the evaluation uses a simulated benchmark rather than production traces from a deployed enterprise agentic system; while the simulated task chains were constructed to reflect realistic tool-calling patterns in the three target domains, absolute metric values should be treated as illustrative rather than as generalizable production estimates, and the framework should be replicated against real, de-identified enterprise agent logs as a priority next step. Second, the composite Agentic Risk Index depends on domain-calibrated weighting coefficients (α , β , γ) and a carry-forward parameter (λ) whose optimal values were set by illustrative grid search rather than by a formally validated calibration procedure; future work should develop a principled calibration methodology, potentially incorporating expert elicitation alongside outcome-based tuning. Third, the framework as presented treats tool-call schema and scope validation as deterministic gates; extending these gates to reason probabilistically about ambiguous or partially specified tool calls is an open technical problem. Fourth, this study evaluates a bounded task-chain length (six to twenty steps); longer-horizon agentic workflows, and multi-agent systems in which several agents interact, introduce additional compounding and coordination risks not directly modeled here and represent a natural extension of this work.

9. Conclusion

This paper has argued that the enterprise adoption of knowledge-grounded and agentic AI systems introduces validation and governance risks — grounding failure, tool-call malformation, permission-scope violation, and compounding chain risk — that are structurally distinct from, and additive to, the hallucination and bias risks addressed by prior enterprise AI assurance work. The proposed Agentic and Knowledge-Grounded Assurance (AKGA) framework operationalizes detection and mitigation of these risks across a six-layer governance pipeline, combining per-statement grounding-fidelity scoring, pre- and post-execution tool-call verification, action-safety scoring, and a carry-forward composite risk index with tiered, threshold-based escalation. Simulated benchmark results across healthcare, financial-services, and insurance domains show substantial post-governance improvement in all component metrics and demonstrate that step-wise verification meaningfully suppresses the compounding of risk across multi-step agentic task chains relative to an ungoverned baseline. As enterprises move from generative assistance toward autonomous, tool-using agentic execution, the central claim of this paper is that

validation and governance of such systems must be engineered as a continuous, auditable property of the

executing agent's full trace — not inferred from the quality of its final output alone.

References

- Lewis, P., Perez, E., Piktus, A., Petroni, F., Karpukhin, V., Goyal, N., Küttler, H., Lewis, M., Yih, W., Rocktäschel, T., Riedel, S., & Kiela, D. (2020). Retrieval-augmented generation for knowledge-intensive NLP tasks. *Advances in Neural Information Processing Systems*, 33, 9459–9474.
- Yao, S., Zhao, J., Yu, D., Du, N., Shafran, I., Narasimhan, K., & Cao, Y. (2023). ReAct: Synergizing reasoning and acting in language models. *Proceedings of the 11th International Conference on Learning Representations (ICLR)*.
- Schick, T., Dwivedi-Yu, J., Dessì, R., Raileanu, R., Lomeli, M., Zettlemoyer, L., Cancedda, N., & Scialom, T. (2023). Toolformer: Language models can teach themselves to use tools. *arXiv:2302.04761*.
- Shinn, N., Cassano, F., Berman, E., Gopinath, A., Narasimhan, K., & Yao, S. (2023). Reflexion: Language agents with verbal reinforcement learning. *arXiv:2303.11366*.
- Park, J. S., O'Brien, J., Cai, C. J., Morris, M. R., Liang, P., & Bernstein, M. S. (2023). Generative agents: Interactive simulacra of human behavior. *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology*.
- Nakano, R., Hilton, J., Balaji, S., Wu, J., Ouyang, L., Kim, C., Hesse, C., Jain, S., Kosaraju, V., Saunders, W., et al. (2021). WebGPT: Browser-assisted question-answering with human feedback. *arXiv:2112.09332*.
- Ruan, Y., Dong, H., Wang, A., Pitis, S., Zhou, Y., Ba, J., Dubois, Y., Maddison, C. J., & Hashimoto, T. B. (2024). Identifying the risks of LM agents with an LM-emulated sandbox. *Proceedings of the 12th International Conference on Learning Representations (ICLR)*.
- Ji, Z., Lee, N., Frieske, R., Yu, T., Su, D., Xu, Y., Ishii, E., Bang, Y. J., Madotto, A., & Fung, P. (2023). Survey of hallucination in natural language generation. *ACM Computing Surveys*, 55(12), 1–38.
- Wei, J., Wang, X., Schuurmans, D., Bosma, M., Ichter, B., Xia, F., Chi, E., Le, Q., & Zhou, D. (2022). Chain-of-thought prompting elicits reasoning in large language models. *Advances in Neural Information Processing Systems*, 35, 24824–24837.
- National Institute of Standards and Technology (NIST). (2023). Artificial Intelligence Risk Management Framework (AI RMF 1.0) (NIST AI 100-1). U.S. Department of Commerce. <https://doi.org/10.6028/NIST.AI.100-1>
- European Parliament. (2024). Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). *Official Journal of the European Union*.
- Mitchell, M., Wu, S., Zaldivar, A., Barnes, P., Vasserman, L., Hutchinson, B., Spitzer, E., Raji, I. D., & Gebru, T. (2019). Model cards for model reporting. *Proceedings of the ACM Conference on Fairness, Accountability, and Transparency*, 220–229.
- Kamiran, F., & Calders, T. (2012). Data preprocessing techniques for classification without discrimination. *Knowledge and Information Systems*, 33(1), 1–33.
- Hardt, M., Price, E., & Srebro, N. (2016). Equality of opportunity in supervised learning. *Advances in Neural Information Processing Systems*, 29, 3315–3323.
- Varshney, N., Yao, W., Zhang, H., Chen, J., & Yu, D. (2023). A stitch in time saves nine: Detecting and mitigating hallucinations of LLMs by validating low-confidence generation. *arXiv:2307.03987*.
- Chan, A., Salganik, R., Markelius, A., Pang, C., Rajkumar, N., Krashennnikov, D., Langosco, L., He, Z., Duan, Y., Carroll, M., et al. (2023). Harms from increasingly agentic algorithmic systems. *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency*, 651–666.