

CognitiveInsight LLC

AI-Enabled GRC and Internal Control Assurance

A Governance Architecture for Agentic Workflows, Financial-Control Integrity, and Verifiable Evidence

Includes Cryptographic Evidence Receipts & Verifiable Provenance Addendum

Author's Governance Principle

When AI participates in a control, audit, or financially consequential workflow, the business process and the AI system must be assured independently. AI may assist; accountable controls must authorize.

James Greenwood

Founder, CognitiveInsight LLC

Version 1.1 - Architectural and Governance Position Paper - August 2026

A Note on Authorship and AI Assistance

This paper is my own work. The governing concept - that an AI-enabled process must not become self-validating, and that the business process, the AI system, and the governance function around it require independent assurance - is a position I formed through my own research and professional experience connecting GRC, internal-control, and AI-governance practice. The frameworks original to this paper, including the three-layer assurance architecture, the control heuristic, the Operational Authority Matrix, and the assurance lifecycle, are my synthesis, not restatements of any single external standard.

I used an AI assistant to help draft, structure, and edit this paper - organizing sections, tightening language, building the tables, and checking that cited external sources say what I intended to attribute to them. I reviewed the result and take responsibility for the final content.

List of AI used with purpose:

- Claude, by Anthropic - Honest evaluation of citation and usage in the content of the text.
- Google Gemini Deep research – Compilation of conceptual thoughts as related to internal audit, accounting, GRC, Automated deterministic process used for audit
- Perplexity to refine the citations and wording

Note: Each model is trained on a unique set of data and that has a relationship to the best use of each models resources and abilities.

CognitiveInsight LLC is my own entity. It is also the name behind cognitiveinsight.ai, where I publish the Cognitive Insight Audit Framework (CIAF) and Lazy Capsule Materialization (LCM) - an independent, non-peer-reviewed research initiative I created for cryptographically verifiable AI audit trails. Where the Addendum to this paper references that work, it is citing my own related research, not an independent third party. I am noting that plainly so a reader can weigh it accordingly: CIAF/LCM is one illustrative implementation pattern I am proposing, not an adopted standard, a regulatory framework, or third-party validation of the governance model described here.

Everything else cited in this paper - NIST's AI Risk Management Framework, the EU AI Act, ISO/IEC 42001, the OWASP Top 10 for LLM Applications, AICPA/CIMA's SOC suite, and the PCAOB's 2024 staff Spotlight on generative AI in audits - is an independent, external, publicly available source. Citations are attached only where they support the specific point being made, rather than to every sentence. Where a claim is my own framework or judgment rather than something drawn from one of these sources, it is labeled "Author's Framework" instead of implying external validation it doesn't have.

My goal in publishing this is to help people - GRC leaders, auditors, risk and compliance teams, and AI system owners - reason clearly about a genuinely complicated problem: how to get value from AI in control and audit workflows without quietly letting AI become the thing that grades its own homework. If this framing is useful to your organization, please use it, adapt it, and challenge it.

Executive Summary

Organizations are evaluating AI for a growing range of operational, risk, and assurance workflows. It now extracts evidence, maps policies, triages exceptions, drafts risk assessments, summarizes contracts, monitors controls, and may influence financial reporting, regulatory disclosures, access decisions, and customer outcomes. This creates a governance problem that traditional automation programs and high-level responsible-AI principles do not fully resolve on their own.

The central proposition of this paper is simple: an AI-enabled process must not become self-validating. When an AI system helps test a control, interpret evidence, or recommend a consequential decision, the organization must assure both (1) the underlying business control and (2) the AI-enabled workflow that influenced the conclusion.

- **AI Can Assist ≠ AI Can Authorize:** AI may accelerate interpretation, extraction, retrieval, synthesis, and anomaly triage; it should not become the unverified final authority for material conclusions or irreversible actions.
- **Primary Evidence Authority:** Primary evidence remains the authoritative basis for a control conclusion; AI-generated outputs are derived interpretations that must remain traceable to that evidence.
- **Proportional Control Rigor:** Risk-based controls should increase with materiality, autonomy, privilege, and irreversibility - especially where AI affects financial reporting, key controls, regulated decisions, or production systems.
- **Comprehensive Assurance Baseline:** Trustworthy assurance requires provenance, deterministic validation, constrained permissions, lifecycle evaluation, independent testing, monitoring, and accountable ownership.

Scope and Use

This paper provides an enterprise architecture and governance model. It is not legal advice, an audit opinion, or a determination that a particular use case is subject to a specific law, assurance engagement, or reporting requirement. Organizations should apply the model with legal counsel, finance, security, privacy, risk, internal audit, and external-assurance stakeholders as appropriate.

1. The Assurance Problem

Rule-based automation can provide highly reproducible results when its inputs, rules, dependencies, execution environment, and state are controlled. AI-enabled systems introduce additional uncertainty through probabilistic inference, retrieval context, dynamic prompts, tool selection, external data, model updates, and human reliance on generated outputs. The risk is not that AI is inherently unusable; the risk is that it can be given authority without commensurate evidence, constraints, and oversight.

Avoiding Circular Assurance

Author's Governance Principle

An AI agent that identifies a control failure - or declares a control effective - cannot be the sole proof of its own conclusion. The organization must preserve an independent path from source evidence to final decision.

Assurance Target	Question to Answer
Business Process or Control	Is the underlying requirement met? Are transactions, records, configurations, estimates, or activities complete, accurate, authorized, and timely?
AI-Enabled Workflow	Can the AI system securely and reliably perform its approved role using the right evidence, logic, permissions, and escalation rules?
Governance & Assurance Model	Are ownership, risk acceptance, independent testing, change control, monitoring, and remediation operating effectively?

Recommended Architectural Pattern

AUTHOR'S FRAMEWORK - ORIGINAL TO THIS PAPER

Primary Evidence → AI Interpretation → Deterministic Validation → Authorized Disposition → Versioned Audit Record

This pattern separates fact, interpretation, validation, and decision. It permits AI to create value while preserving a defensible basis for assurance.

2. A Layered Assurance Architecture

Mature organizations should treat AI-enabled assurance as three connected - but distinct - layers. A failure at one layer should be detectable and containable by another.

Layer	Purpose	Representative Controls
1. Business-Process Assurance	Assure the transactions, controls, financial assertions, security requirements, and operational obligations under review.	Control objectives; reconciliations; approvals; policy tests; source-system records; exception management; financial-review controls.
2. AI-System Assurance	Assure the AI model, prompts, retrieval, tooling, permissions, outputs, and failure handling that influence the process.	Model and agent inventory; evaluations; grounding; tool allowlists; data controls; red teaming; logging; drift detection; change review.
3. Independent Governance Assurance	Assure that ownership, risk decisions, monitoring, internal audit, and remediation are independent and effective.	Board/audit-committee reporting; risk acceptance; control testing; internal audit; third-party assurance; incident exercises; remediation SLAs.

Operational Definition of Independent Assurance

For purposes of this paper, independence means that the function assessing the AI-enabled workflow has sufficient authority, evidence access, relevant technical competence, and organizational separation from the team that builds or operates the workflow to challenge its design, controls, and reported performance.

Responsibility Model

- **Business Owner:** Accountable for the business outcome, process design, and use of AI within the approved purpose.
- **AI System Owner:** Accountable for the model, prompts, orchestration, data connections, evaluations, monitoring, and operational remediation.
- **Risk and Compliance:** Defines use-case classification, policy obligations, control objectives, evidence requirements, exception thresholds, and risk-acceptance processes.
- **Security and Privacy:** Assess identity, access, data protection, prompt-injection exposure, external connections, logging, and incident response.
- **Finance Controllorship:** Governs accounting policy relevance, review requirements, and evidence supporting estimates, entries, reconciliations, and disclosures.
- **Internal Audit:** Independently assesses the design and operating effectiveness of the AI governance and control environment.
- **Board or Audit Committee:** Provides oversight for material AI risk, financial-reporting implications, significant incidents, control deficiencies, and management remediation.

3. AI in the Internal-Control Environment

AI governance establishes enterprise rules for responsible AI. Internal control extends those rules into critical business processes. If an AI output can influence a material transaction, accounting estimate, journal entry, reconciliation, financial disclosure, management-review control, regulated outcome, or supporting evidence, the AI-enabled workflow should be assessed as part of the internal-control environment.

The use of AI does not eliminate management's responsibility to apply the applicable accounting framework, maintain adequate supporting evidence, and design effective controls over the affected process. For organizations subject to financial-reporting control requirements, the relevant question is whether AI use affects the design or operation of controls over financial reporting - not whether the system is branded as "AI."

AI-Enabled Activity	Primary Risk	Control Posture
Drafting non-financial content	Inaccuracy or confidential-data exposure	Approved-tool policy, data restrictions, qualified accountable review.
Summarizing contracts or policies	Omitted terms or incorrect interpretation	Citations to source clauses, exception routing, qualified legal/finance review.
Transaction anomaly detection	False positives, false negatives, biased prioritization	Back-testing, governed thresholds, independent investigation and disposition.

AI-Enabled Activity	Primary Risk	Control Posture
Reconciliations & exception matching	Incomplete matching or incorrect exception resolution	Read-only access, deterministic reconciliation logic, reviewer sign-off, retained evidence.
Journal-entry drafting	Unsupported, inaccurate, or unauthorized entry	For material entries: prohibit autonomous posting unless a formal risk assessment, control design, authorization model, and evidence standard support a limited exception.
Accounting estimates or revenue analysis	Incorrect methodology, assumptions, or framework treatment	Controlled data, approved methodology, sensitivity analysis, qualified review, evidence retention.
Control monitoring or audit support	False pass/fail conclusion or missing evidence	Deterministic evidence tests, model evaluation, exception escalation, independent testing.
Disclosure support	Material omission or misstatement	Verified source data, documented review, finance/legal approval, disclosure-control integration.

Proposed Control Heuristic

AUTHOR'S FRAMEWORK - ORIGINAL TO THIS PAPER

Required Control Rigor $\propto$ Materiality $\times$ Autonomy $\times$ Privilege $\times$ Irreversibility $\times$ Evidence Uncertainty

This formula is a decision heuristic I use to communicate a relationship, not a calibrated quantitative risk calculation, and it is not drawn from NIST, ISO, or any other standard listed in the references - it is original to this paper. Organizations should define thresholds and scoring approaches appropriate to their own enterprise risk management methodology. A read-only agent summarizing public guidance warrants a different control posture from an agent that can post a journal entry, modify identity access, accept a policy exception, or produce a disclosure-supporting calculation. Authority must narrow as potential impact rises.

Operational Authority Matrix

To operationalize “AI may assist; accountable controls must authorize,” organizations should map AI functional capabilities directly to permitted decision rights and required control safeguards. This matrix, like the heuristic above, is my own operational framework.

AI Authority Level	Permitted Operational Role	Required Control Safeguards
Inform	Summarize unstructured text, retrieve policy clauses, classify incoming tickets.	Approved corpus boundaries, source attribution/citations, mandatory qualified review before use.

AI Authority Level	Permitted Operational Role	Required Control Safeguards
Recommend	Propose control findings, draft exception responses, flag transaction anomalies.	Evidence grounding, deterministic schema checks, documented human reviewer disposition.
Execute Reversible Action	Create draft GRC tickets, collect system evidence logs, trigger low-risk workflows.	Scoped API credentials, tool allowlists, tamper-evident or access-controlled logging, automated rollback capability.
Execute Consequential Action	Modify IAM permissions, post journal entries, grant formal risk exceptions, submit regulatory disclosures.	Prohibited without explicit, authenticated human authorization and multi-layer independent testing.

4. AI Governance as Control Infrastructure

Responsible-AI principles are necessary but should be translated into operational controls. For AI systems within the EU AI Act's high-risk regime, human-oversight measures must be designed in accordance with the applicable requirements of Article 14; for other systems, organizations should apply risk-proportionate human-oversight controls. "Transparency" alone is insufficient if the organization cannot reconstruct the source evidence, model configuration, tool calls, reviewer decision, and exception history behind a material output.

Governance Principle	Operational Meaning	Evidence of Operation
Transparency & Explainability	Document purpose, decision role, limitations, data lineage, model/provider, and workflow logic at a level proportionate to risk.	System card; use-case record; model/version register; source citations; decision trace.
Accountability & Human Oversight	Assign named owners, decision rights, reviewer authority, escalation, and risk acceptance.	RACI; approval records; review criteria; exception and override logs.
Fairness & Equity	Assess bias and disparate impact when systems affect people, eligibility, access, prioritization, or outcomes.	Impact assessment; evaluation results; accessibility/review procedures; appeal or recourse process.
Privacy & Security	Protect data, tools, credentials, and users throughout the workflow.	Data classification; access reviews; retention schedule; threat model; penetration/red-team results.
Reliability & Integrity	Test accuracy, completeness, robustness, grounding, and failure behavior for the approved task.	Evaluation suite; regression results; quality thresholds; incident records; monitoring dashboard.
Auditability & Provenance	Retain an evidence chain from source through interpretation to final decision.	Tamper-evident logs; artifact identifiers; reviewer disposition; retention evidence.
Change Governance	Reassess material changes before or shortly after release under controlled monitoring.	Change ticket; impact assessment; revalidation; rollout approval; rollback plan.

5. Assuring Agentic GRC Systems

Tool-using agents create a distinct risk profile because untrusted inputs, model outputs, and connected tools can expand the attack surface and lead to unsafe downstream actions. Prompt instructions and untrusted documents should never be treated as an authorization mechanism. Tool permissions, policy enforcement, and independent validation must govern execution.

Minimum Control Baseline

- **1. Inventory Management:** Maintain an inventory of AI agents, models, prompts, connectors, data sources, business owners, risk tiers, and permitted actions.

- **2. Least Privilege Boundaries:** Constrain authority through least privilege, scoped credentials, tool allowlists, transaction limits, environment segmentation, and default read-only access.
- **3. Untrusted Content Separation:** Separate untrusted content from executable authority. Treat documents, emails, tickets, repository metadata, web content, and tool descriptions as data - not instructions.
- **4. Citation Grounding:** Ground material outputs in approved source artifacts and preserve citation-level evidence for each material claim.
- **5. Deterministic Validation:** Use deterministic validation for structured facts, policy rules, calculations, entitlements, and material control conclusions.
- **6. Task-Specific Pre-Deployment Testing:** Perform task-specific pre-deployment and regression evaluations for accuracy, completeness, policy adherence, prompt injection, tool-call correctness, data leakage, and escalation behavior.
- **7. Human Approval Gates:** Require qualified, accountable human review for irreversible, high-impact, financially material, or rights-affecting actions.
- **8. Comprehensive Telemetry:** Log inputs, source identities, retrieval results, prompt and workflow versions, model/version, settings, tool calls, validation results, reviewer decisions, and timestamps.
- **9. Continuous Oversight:** Operate monitoring, incident response, credential revocation, workflow suspension, and rollback capabilities.

Cryptographic Receipts and Verifiable Provenance

For material AI-assisted decisions, organizations should preserve a verifiable evidence chain connecting authoritative source artifacts, model and workflow configuration, retrieval context, tool invocations, deterministic validation results, and human or policy disposition. Cryptographic receipts can strengthen the integrity and independent verifiability of that chain. They do not establish substantive correctness on their own; they complement source validation, control testing, qualified review, and independent assurance.

Testing the AI That Audits

The agent should be evaluated against task-specific control objectives and failure modes, not merely an aggregate accuracy score; relevant testing includes validity, robustness, source-grounding, security, and real-world operating context. Test cases should include ordinary scenarios, ambiguous evidence, conflicting sources, missing information, adversarial instructions, policy exceptions, model/provider changes, and tool failures. Independent internal audit or an equivalent second-line/third-line function should validate that the testing program itself is designed and operating effectively.

Worked Use Case: AI-Assisted SOC 2 Evidence Review in Third-Party Risk Management

To demonstrate how this three-layer assurance architecture and control baseline operate in practice, consider an automated Third-Party Risk Management (TPRM) workflow evaluating vendor compliance:

- **Step 1 - Ingestion & Source Hash:** A third-party vendor submits a SOC 2 Type II report via the procurement portal. The workflow engine calculates a SHA-256 hash of the submitted PDF and retains the document, or an authorized evidence reference, in tamper-evident, retention-governed, access-controlled storage.
- **Step 2 - AI Extraction & Grounding:** An LLM agent parses the unstructured PDF to extract control statements, report period, auditor opinion, test scope, and Complementary User Entity Controls (CUECs). Every extracted field retains citation metadata linking back to page and paragraph numbers.
- **Step 3 - Deterministic Policy Engine Validation:** Extracted data flows into a deterministic policy engine (e.g., Open Policy Agent). Hard rules verify against organizational vendor-risk criteria: does the report satisfy the approved recency threshold; does the auditor's opinion meet policy requirements after assessing any modification, scope limitation, or noted exception; does the report cover the required Trust Services Criteria. If any rule fails, the status flags as “Non-Compliant” or “Escalated.”
- **Step 4 - Qualified Accountable Human Review:** If validation flags CUEC gaps or modified conclusions, the workflow routes the package to an authorized vendor risk analyst with the required decision rights, who reviews grounded citations and records a formal, authenticated disposition.
- **Step 5 - Cryptographic Evidence Receipt Generation:** Upon authorization, the system generates a receipt binding the source document hash, extraction schema, prompt/model version, deterministic check results, and reviewer signature.
- **Step 6 - Recommended Assurance Procedure:** Internal or external auditors should sample completed vendor reviews, verifying receipt signatures against source PDFs and confirming that no policy bypasses occurred.

6. Lifecycle for AI Assurance

Stage	Required Activity	Key Output
1. Intake & Classification	Define purpose, users, affected parties, data, tools, autonomy, financial relevance, and regulatory context.	Use-case record and risk tier.
2. Control Design	Set authority boundaries, data restrictions, validations, review gates, monitoring, retention, and incident procedures.	Control design and RACI.
3. Evaluation & Approval	Test quality, grounding, security, fairness where relevant, tool behavior, and failure modes.	Evaluation report and release decision.
4. Deployment	Implement controlled access, versioning, telemetry, user guidance, and kill-switch mechanisms.	Production baseline and evidence package.

Stage	Required Activity	Key Output
5. Continuous Monitoring	Detect quality degradation, anomalous access, unsupported outputs, prompt injection, policy violations, and overrides.	Metrics, alerts, and periodic review.
6. Change Management	Assess model, prompt, data, tool, policy, connector, and workflow changes.	Impact analysis and revalidation.
7. Independent Assurance	Test control design and operating effectiveness; report deficiencies and remediation.	Internal-audit or assurance results.
8. Retirement	Revoke access, preserve required records, migrate dependencies, and formally close risk ownership.	Decommissioning evidence.

7. Executive Metrics and Reporting

Executives and audit committees need decision-useful indicators rather than generic AI-usage counts. Reporting should show where AI exists, what authority it has, whether it is performing within approved bounds, and whether residual risk is increasing or decreasing.

- **AI Inventory Coverage:** Percentage of production AI use cases recorded, classified, owned, and reviewed.
- **High-Risk Coverage:** Percentage of high-impact use cases with current evaluations, security review, privacy review, and approved control design.
- **Evidence-Grounding Rate:** Percentage of material outputs linked to authoritative source artifacts.
- **Validation Exception Rate:** Outputs rejected, escalated, or corrected by deterministic checks or reviewers.
- **Unauthorized Tool-Call Rate:** Attempted actions outside approved scope, including anomalous access patterns.
- **Human-Override Rate:** Disagreement rate useful for detecting quality degradation, reviewer burden, or unclear operating criteria.
- **Change Recertification Timeliness:** Percentage of material changes assessed and revalidated before production release.
- **Open Risk Aging:** Material control gaps, owners, due dates, compensating controls, and risk-acceptance status.

8. Implementation Roadmap

- **First 30 Days - Establish Visibility:** Create the AI-use-case inventory; identify agents connected to sensitive data, financial workflows, production systems, or material controls; assign accountable owners; immediately review high-privilege tool access.

- **Days 31–60 - Set Control Baseline:** Classify use cases; define risk tiers; implement minimum logging, access, evidence, and review requirements; establish a change-management trigger for model and workflow changes.
- **Days 61–90 - Evaluate & Assure:** Deploy task-specific evaluation suites; red-team untrusted-input and tool-use pathways; formalize independent testing with internal audit and control owners; report the first executive risk dashboard.
- **Ongoing - Continuous Assurance:** Integrate AI governance with enterprise risk, internal control, security engineering, privacy, vendor risk, financial controllership, and audit planning; tune controls using incidents, overrides, evaluation trends, and changing regulatory expectations.

Addendum: Cryptographic Evidence Receipts and Verifiable Provenance

Technical Addendum | August 2026

This addendum extends the core AI-assurance architecture with a technical pattern for independently verifiable evidence lineage in material AI-assisted workflows.

Addendum Thesis

When AI contributes to audit, GRC, internal-control, financial-reporting, access, or other consequential workflows, ordinary application logs may not provide a sufficiently independent or tamper-evident evidence trail. Organizations should consider cryptographically verifiable receipts that bind the relevant source evidence, system configuration, validation outcomes, and authorized disposition into an evidence chain that an authorized independent reviewer can test.

Author's Provenance Principle

AUTHOR'S FRAMEWORK - ORIGINAL TO THIS PAPER

Evidence Integrity Does Not Equal Evidence Correctness

Cryptographic provenance can help establish that a retained record is attributable and resistant to undetected modification within a defined capture boundary. It does not establish that the source evidence was truthful, complete, or properly scoped; that the AI model interpreted it correctly; or that an accounting, compliance, or operational conclusion was substantively correct. Those matters still require evidence validation, policy controls, deterministic checks where suitable, qualified review, and independent assurance.

Cognitive Insight as an Implementation Pattern

Disclosure

Cognitive Insight (cognitiveinsight.ai) is my own independent research initiative, published under the same authorship as this paper. It is not a third-party validation of this framework, an accredited standard, or a regulator.

Cognitive Insight describes the Cognitive Insight Audit Framework (CIAF) and Lazy Capsule Materialization (LCM) as an independent research initiative for cryptographically verifiable AI governance and auditability. The initiative proposes mechanisms involving verifiable receipts, integrity proofs, and verification workflows. In this paper's architecture, it is best used as a potential implementation pattern for provenance and evidence integrity - not as a replacement for regulatory interpretation, accounting judgment, internal audit, or formal external assurance.

The initiative proposes alignment with the EU AI Act, NIST AI RMF, and ISO/IEC 42001. Actual regulatory or statutory compliance depends on the specific enterprise context, implementation scope, and authorized legal/audit evaluation - alignment with these frameworks is a design goal of CIAF/LCM, not a certification.

What a Material Receipt Should Bind

- **Workflow Identity:** AI workflow identity, approved use case, risk tier, business owner, and system owner.
- **Model & Configuration Baseline:** Model provider, model/version identifier, inference settings, prompt-template version, workflow/orchestration version, and policy version.
- **Evidence References:** Identity, version, hash, and retention reference for authoritative source evidence and retrieved context. Raw content should be retained only when authorized and necessary.
- **Tool Interactions:** Tool and connector identifiers, authorization scope, requested action, system response, and resulting state where applicable.
- **Validation Outcomes:** Deterministic validation rules, input/output hashes, validation results, exceptions, and escalation status.
- **Human/Policy Disposition:** Authorized human or policy disposition, reviewer identity, decision rationale, override reason, and timestamp.
- **Independent Proofs:** Ordered timestamps, integrity proof, and a verification procedure available to authorized independent reviewers.

Where Provenance Fits

Assurance Layer	Contribution of Verifiable Provenance	What It Does Not Replace
Business-Process Assurance	Links material conclusions to traceable source artifacts, validations, and approval records.	Accounting policy, control design, substantive testing, reconciliations, or qualified review.
AI-System Assurance	Preserves model/workflow versions, retrieval context, tool activity, validator results, and output lineage.	Task-specific evaluation, red teaming, access control, monitoring, or quality thresholds.
Independent Governance Assurance	Supports independent verification without sole reliance on an operator dashboard or mutable application database.	Professional skepticism, audit methodology, management accountability, or risk acceptance.

AUTHOR'S FRAMEWORK - ORIGINAL TO THIS PAPER

***Authoritative Evidence → AI Interpretation → Deterministic Validation → Authorized Disposition
→ Cryptographic Receipt → Independent Verification***

The receipt is not the truth of a business event. It is an integrity-preserving record of the controlled path through which the organization formed, validated, and authorized a conclusion.

Control Design Considerations for Provenance

- **1. Capture Boundaries:** Define the capture boundary. State precisely which events, systems, artifacts, exceptions, and actions are covered; do not imply completeness outside that scope.

- **2. Confidentiality & Hashing:** Preserve confidentiality. Use hashes, encrypted evidence, and access-controlled references where raw prompts, documents, output, or personal data should not be broadly retained.
- **3. Segregation of Duties:** Separate duties. The workflow operator should not be the only party able to alter proof infrastructure, verification logic, retention settings, or access rights.
- **4. Independent Sampling:** Test independently. Periodically sample receipts and reconcile them to authoritative source systems, retained artifacts, workflow records, and final dispositions.
- **5. Records Governance:** Integrate records governance. The implementation must support records retention, privacy deletion requirements, legal holds, contractual restrictions, and incident investigation needs.
- **6. Controlled Change Triggers:** Treat material changes as controlled events. Changes to the model, prompt, tool, connector, policy, source corpus, or workflow should trigger impact analysis and revalidation.

Positioning and Citation Note

Cognitive Insight should be cited as my own independent research initiative for statements concerning the CIAF/LCM framework, technical specifications, demonstrations, and proposed methods. It should not be represented as an official regulator, accredited assurance standard, or automatic proof of compliance with the EU AI Act, GAAP, SOX, SOC examinations, ISO/IEC 42001, or other frameworks. Claims of compliance with any of those frameworks require the relevant legal, professional, or standards-body authority - not this paper or CIAF/LCM.

Conclusion

AI can make GRC and internal control more continuous, evidence-rich, and responsive. It can reduce manual effort in document processing, evidence retrieval, policy mapping, anomaly detection, and triage. Those benefits are real - but they do not eliminate the need for control design. They increase it.

The durable enterprise model is not autonomous compliance. It is bounded intelligence operating inside accountable controls. Organizations should ensure that primary evidence remains authoritative, material claims are validated, consequential actions are authorized, AI systems are independently assured, and governance remains capable of stopping or correcting the system when it fails.

Selected Standards, Regulatory Sources, and Research References

External, Independent Sources

- 1. NIST, AI Risk Management Framework (AI RMF 1.0), NIST AI 100-1 (2023). nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-1.pdf
- 2. NIST, Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile, NIST AI 600-1 (2024). nvlpubs.nist.gov/nistpubs/ai/NIST.AI.600-1.pdf
- 3. NIST, AI Risk Management Framework resource center. nist.gov/itl/ai-risk-management-framework
- 4. OWASP GenAI Security Project, Top 10 for Large Language Model Applications (accessed August 2026). owasp.org/www-project-top-10-for-large-language-model-applications/
- 5. Regulation (EU) 2024/1689 (EU Artificial Intelligence Act), EUR-Lex. eur-lex.europa.eu/eli/reg/2024/1689/oj
- 6. ISO/IEC 42001:2023, Artificial intelligence - Management system. iso.org/standard/81230.html
- 7. AICPA & CIMA, System and Organization Controls (SOC) suite of services. aicpa-cima.com/resources/landing/system-and-organization-controls-soc-suite-of-services
- 8. PCAOB, Staff Update on Outreach Activities Related to the Integration of Generative Artificial Intelligence in Audits and Financial Reporting (Spotlight, July 22, 2024). pcaobus.org

Author's Own Related Research (Not an Independent Third Party)

- 9. Greenwood, J., Cognitive Insight - AI Governance Research & CIAF/LCM Framework (accessed August 2026). cognitiveinsight.ai - published by the same author and entity as this paper; cited here as an illustrative implementation pattern, not as independent corroboration.