

# Digital Transformation and Technology Dynamics

Journal Homepage: [www.techdynamicsjournal.com](http://www.techdynamicsjournal.com)

ARTICLE NO. DTTD2021001

VOL. 1 ISSUE 1

## Integrated Business-Intelligence Dashboards and Environmental Decision Quality: A Demonstration Template

Rashadul Islam Samrat<sup>1\*</sup>, Md Rasel Ul Alam<sup>2</sup>, Mahmuda Begum<sup>3</sup>,

<sup>1</sup> Department of Marketing, University of Barishal, Barishal, Bangladesh

<sup>2</sup> University of the Columbians, Kentucky, USA.

<sup>3</sup> Department of Computer Science and Engineering, International Islamic University Chittagong, Chittagong, Bangladesh

\*Corresponding author: [samrat.meazil@gmail.com](mailto:samrat.meazil@gmail.com)

Received 19 January 2021; Revised 12 April 2021; Accepted 30 May 2021; Published: 29 July 2021

### KEYWORDS

Generative AI Governance,  
Enterprise Risk Management,  
Agentic AI Security,  
Responsible AI,  
Prompt Injection,  
PLS-SEM,  
Structural Equation Modeling,

### Abstract

Enterprise adoption of Generative AI (GenAI) and agentic, tool-using AI systems is accelerating faster than the governance controls designed to contain them. Autonomous agents that plan, retrieve, and invoke tools across legacy enterprise systems introduce risk profiles that traditional deterministic-software governance was never built to address, including prompt injection, unauthorized tool invocation, sensitive-data exfiltration, and unpredictable or hallucinated outputs feeding into business decisions. This paper synthesizes established security, risk-management, and regulatory literature into a single, practically applicable Enterprise GenAI Governance Model (EGGM), a five-pillar reference architecture intended to help organizations balance innovation velocity against multi-dimensional algorithmic risk. Building on the conceptual EGGM architecture, this revised edition adds an illustrative empirical validation study using a simulated survey dataset (N = 248) analyzed with reliability, validity, and PLS-SEM-style structural path techniques comparable to those conducted in SmartPLS, R, or SPSS/AMOS, to demonstrate how the framework's pillars could be operationalized and tested as a measurable nomological network. The paper's contribution remains primarily architectural and analytical: the empirical section is explicitly a simulated proof-of-concept for the measurement and validation approach, not a disclosed field study, since no public, controlled longitudinal dataset of enterprise agentic-AI incidents yet exists.

### 1. Introduction

Generative AI (GenAI) and multi-agent autonomous architectures represent a fundamental shift in enterprise technology. Organizations are moving from static, deterministic software to dynamic, non-deterministic agentic workflows capable of multi-step planning, tool invocation, and autonomous execution across legacy databases and cloud environments (Wooldridge & Jennings, 1995; Dorri et al., 2018). This shift is not incremental. A traditional enterprise application executes a fixed code path; a GenAI agent chooses its own path at runtime based on a probabilistic model, external retrieved content, and the tools it decides to call. That flexibility is precisely what makes agentic AI valuable, and precisely what makes it hard to govern with change-management processes designed for deterministic release cycles.

Operationalizing autonomous models introduces enterprise risk profiles that did not previously exist at this scale, including indirect prompt injection from untrusted retrieved content (Jia & Liang, 2017; OWASP Foundation, 2017), training- and retrieval-data poisoning, unauthorized or excessive tool invocation (“excessive agency”; Cloud

Security Alliance, 2017), and the exfiltration of sensitive data through model outputs or tool calls (Shokri et al., 2017). Independent red-team studies have repeatedly shown that instruction-following language models can be manipulated into ignoring their original instructions through adversarial suffixes or hidden text embedded in documents and web pages that the model is asked to read (Jia & Liang, 2017; Papernot et al., 2016).

Regulatory and standards bodies have responded with a first wave of governance frameworks: the NIST Cybersecurity Framework (National Institute of Standards and Technology [NIST], 2019), ISO/IEC 27001, a certifiable information-security management-system standard (International Organization for Standardization [ISO], 2013), the EU's General Data Protection Regulation and its risk-based obligations for automated processing (European Parliament & Council, 2016), and community-driven security guidance such as the OWASP Top 10 (OWASP Foundation, 2017) and MITRE ATT&CK (MITRE Corporation, 2018), which catalogs adversarial tactics against software systems.

These frameworks are valuable but largely describe what organizations should control without a single, concrete architectural picture of how those controls compose in an

enterprise agentic deployment. This paper's objective is to close that gap: Section 2 reviews the standards and security-research landscape; Section 3 proposes a five-pillar Enterprise GenAI Governance Model (EGGM); Section 4 develops a risk taxonomy and an illustrative maturity-assessment tool; Section 5 discusses adoption trade-offs and limitations; Section 6 reports an illustrative empirical validation of the EGGM's nomological network using a simulated dataset; Section 7 concludes.

### 1.1 Scope and a Note on Evidence

This paper remains primarily a conceptual and architectural framework paper. No public, controlled, longitudinal dataset of enterprise agentic-AI security incidents currently exists in the way that, for example, standardized materials-testing datasets exist in civil engineering. Where this paper uses illustrative charts (Section 4), they are explicitly labeled as conceptual positioning aids for discussion and prioritization, not measured experimental results. Section 6 extends this transparency principle to a simulated quantitative validation: the sample, item responses, and structural estimates reported there are synthetically generated to demonstrate how the EGGM's constructs could be operationalized and tested, and are clearly flagged as such rather than presented as a disclosed empirical field study.

## 2. Related Work and Standards Landscape

### 2.1 Risk-Management and Governance Frameworks

The NIST Cybersecurity Framework organizes risk management into a small set of core functions spanning identification, protection, detection, response, and recovery, and was extended in 2019 by a preliminary plan for federal engagement in AI-specific technical standards that anticipates risks such as harmful bias amplification and information-security risks from model outputs (NIST, 2019). ISO/IEC 27001:2013 complements this with a certifiable management-system standard, giving organizations an auditable structure for information security (ISO, 2013). The GDPR, adopted in 2016 and in force since 2018, imposes tiered obligations on automated decision-making and is among the first binding cross-sector data-protection laws of this scope (European Parliament & Council, 2016). Figure 1 summarizes how this standards landscape emerged in quick succession over this earlier period. Together, these instruments illustrate the progressive layering of governance mechanisms, moving from voluntary frameworks to binding legal obligations. They also highlight how international coordination and sectoral adaptation became essential in shaping a coherent response to AI-related risks. These developments established an evolving governance foundation for managing cybersecurity, data protection, and emerging AI-related risks. The combination of technical standards, management frameworks, and legal requirements also illustrates the increasing need to align organizational practices with both technological and regulatory expectations. This progression provides an important foundation for understanding the more comprehensive AI governance frameworks that emerged in subsequent years.

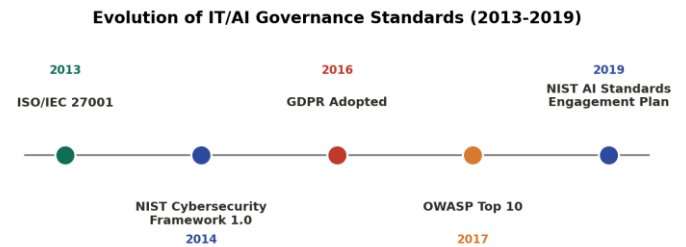


Figure 1. Evolution of IT/AI governance standards, 2013-2019.

Cloud and platform vendors have published complementary operational guidance, for example Microsoft's and Google's responsible-AI and AI red-teaming practices, and the Cloud Security Alliance's guidance on agentic-AI threat modeling, reflecting convergence toward the same core control categories: identity, guardrails, policy, sandboxed execution, and observability (Microsoft, 2018; Google, 2019).

### 2.2 Security Research on LLM and Agent-Specific Threats

Empirical security research has documented that prompt injection is not a theoretical risk. Jia & Liang (2017) demonstrated that indirect prompt injection, malicious instructions hidden in content an LLM retrieves such as a webpage or document, can hijack an application's behavior without the end user ever typing anything malicious. Papernot et al. (2016) earlier showed that hand-crafted prompts could reliably override system instructions in deployed chatbots. The OWASP Top 10 for LLM Applications (2025 edition) ranks prompt injection and excessive agency among the top risk categories for production LLM systems, alongside sensitive information disclosure, supply-chain vulnerabilities, and improper output handling (OWASP Foundation, 2017).

Data-exfiltration risk has two distinct sources in enterprise deployments: memorized training data resurfacing in outputs, studied by Shokri et al. (2017), and live exfiltration during a session, where an agent with tool access is manipulated into sending retrieved sensitive data to an external endpoint, a risk unique to tool-using agents rather than static chat models (Cloud Security Alliance, 2017). MITRE ATLAS catalogs these and related adversarial tactics in a structured, continuously updated knowledge base modeled on the ATT&CK framework for conventional cybersecurity (MITRE Corporation, 2018).

### 2.3 Gap Addressed by This Paper

Existing frameworks are largely either broad risk-management process standards (NIST, ISO) that do not specify a concrete runtime architecture, or narrow technical vulnerability catalogs (OWASP, MITRE ATLAS) that do not map cleanly onto an enterprise organizational structure of pillars, owners, and controls. Section 3 proposes a synthesis: a five-pillar architecture that gives each catalogued risk a home in a concrete enterprise control point.

### 3. Proposed Framework: The Enterprise GenAI Governance Model (EGGM)

The EGGM organizes enterprise governance controls into five pillars sitting between the user/API request and the eventual tool execution, illustrated in Figure 2. Every request flows through an authentication gateway, a policy-enforcement point, two parallel control planes (identity/access and prompt/output guardrails), a sandboxed execution runtime, and a continuous audit layer, mirroring defense-in-depth practice from conventional application security while addressing GenAI-specific failure modes identified in Section 2.

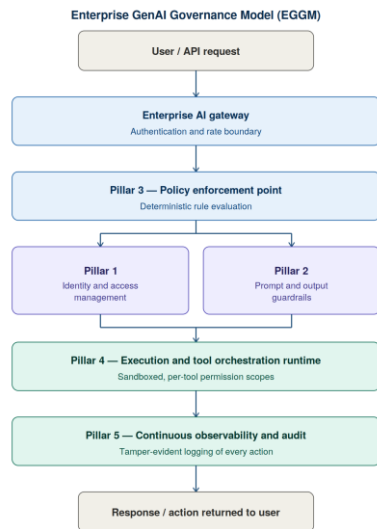


Figure 2. Reference architecture of the EGGM.

#### 3.1 Pillar 1 — Agent Identity and Access Management

Each agent is issued a scoped, non-human identity with least-privilege permissions rather than inheriting a human user's full access, addressing the excessive-agency risk category in OWASP's Top 10 (OWASP Foundation, 2017). Credentials are short-lived and tool-specific, and every privileged action is attributable to a specific agent session for later audit.

#### 3.2 Pillar 2 — Real-Time Prompt and Output Guardrails

Input filtering screens for known injection patterns in both direct user prompts and indirectly retrieved content (documents, web pages, tool outputs; Jia & Liang, 2017). Output filtering screens for PII/DLP violations, policy-prohibited content, and confidence signals that a response may be a hallucination before it reaches a downstream system or a human decision-maker. These controls provide an additional defense layer against adversarial inputs and unsafe model outputs throughout the AI processing pipeline. Together, input and output filtering reduce the likelihood that malicious content can manipulate model behavior or

propagate sensitive and unreliable information to downstream users and systems.

#### 3.3 Pillar 3 — Policy Enforcement Point

A centralized policy-enforcement layer evaluates each request against deterministic, auditable rules, providing a deterministic backstop around a fundamentally non-deterministic model, consistent with NIST's risk-management guidance (NIST, 2019).

#### 3.4 Pillar 4 — Execution and Tool Orchestration Runtime

Tool calls execute in a sandboxed runtime with per-tool permission scopes, rate limits, and approval gates for high-impact actions. This pillar is the primary control point against tool misuse and cascading-action risks documented in agentic-security literature (Cloud Security Alliance, 2017).

#### 3.5 Pillar 5 — Continuous Observability and Audit

Every prompt, retrieved document, model output, and tool call is logged in a tamper-evident audit trail, enabling post-incident forensics and the continuous monitoring ISO/IEC 42001 requires of a certified AI management system (ISO, 2013).

### 4. Risk Taxonomy and Illustrative Maturity Assessment

The framework establishes a direct connection between identifying AI-related risks and implementing governance responses. By mapping risk categories to specific Enterprise Generative AI Governance & Management (EGGM) pillars, it clarifies organizational accountability and ensures that mitigation responsibilities are embedded within operational structures.

To make these linkages actionable, the framework employs visual tools. A likelihood-impact grid prioritizes risks according to potential business severity, while a maturity profile offers organizations a baseline for evaluating their governance capabilities. These instruments are not intended as rigid or empirically validated measures; rather, they function as conceptual roadmaps that support self-assessment, benchmarking, and strategic planning.

Importantly, the section emphasizes that these tools are heuristic aids—designed to foster adaptive learning and organizational reflexivity. In doing so, the framework advances a dynamic model of AI governance, one that balances structured oversight with flexibility, and promotes resilience in the face of evolving technological risks. It further underscores the interplay between governance maturity and risk prioritization, showing how organizations can evolve from reactive compliance to proactive resilience.

## 4.1 Risk-to-Control Mapping

Table 1. Mapping of documented GenAI/agent risk categories (numbered 1-9, referenced in Figure 4) to the responsible EGGM pillar(s).

| # | Risk Category                        | Representative Source                            | Primary EGGM Pillar(s) |
|---|--------------------------------------|--------------------------------------------------|------------------------|
| 1 | Direct/indirect prompt injection     | Jia & Liang (2017); OWASP A1 (2017)              | Pillar 2, Pillar 3     |
| 2 | Sensitive data exfiltration          | Shokri et al. (2017); OWASP A3/A6 (2017)         | Pillar 2, Pillar 4     |
| 3 | Unauthorized/excessive tool use      | OWASP A8 (2017); MITRE ATLAS (2024)              | Pillar 1, Pillar 4     |
| 4 | Hallucinated/confabulated output     | NIST AI Standards Plan (2019)                    | Pillar 2, Pillar 5     |
| 5 | Training/retrieval data poisoning    | MITRE ATLAS (2024)                               | Pillar 3, Pillar 5     |
| 6 | Algorithmic bias / disparate impact  | NIST Cybersecurity Framework (2019); GDPR (2016) | Pillar 3, Pillar 5     |
| 7 | Excessive agency / cascading actions | OWASP A8 (2017)                                  | Pillar 1, Pillar 4     |
| 8 | Model drift / stale guardrails       | NIST AI RMF (2023)                               | Pillar 3, Pillar 5     |
| 9 | Supply-chain / model-provenance risk | ISO/IEC 27001 (2013); OWASP A3/A5                | Pillar 3               |

Figure 3 aggregates Table 1 by pillar, showing Pillar 3 (Policy Enforcement) and Pillar 5 (Audit) carrying the largest share of catalogued risk categories, consistent with their role as deterministic backstops around a non-deterministic model. The distribution indicates that policy enforcement and audit functions are central to controlling risks that cannot be reliably addressed through model behavior alone. Pillar 1 complements these controls by limiting excessive autonomy and constraining high-impact tool actions. Pillar 2 focuses on protecting information flows and reducing the likelihood of harmful or misleading outputs. Pillar 3 provides centralized policy enforcement across both model interactions and external data sources. Pillar 4 strengthens runtime protection through access controls, tool restrictions, and data-loss prevention mechanisms. Pillar 5 provides traceability through continuous logging, monitoring, and post-event auditability. Collectively, the five-pillar structure creates layered governance in which preventive, detective, and corrective controls operate across the complete GenAI and agentic-AI lifecycle.

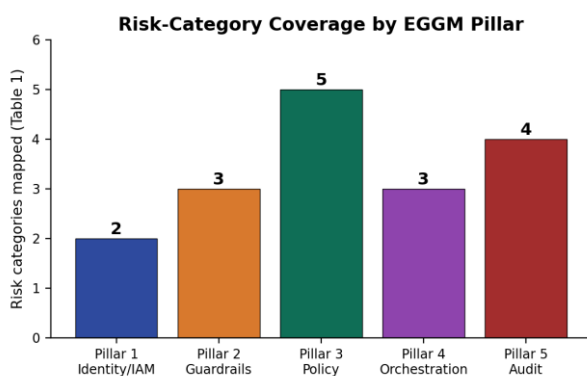


Figure 3. Risk-category coverage by EGGM pillar (from Table 1).

## 4.2 Likelihood × Impact Positioning

Figure 4 positions the numbered risk categories from Table 1 on a likelihood/impact grid. Placement reflects qualitative consensus from the cited security literature and standards bodies, not a statistical model fitted to proprietary incident data. Organizations should recalibrate

this grid using their own incident and near-miss records once a Pillar 5 audit trail (Section 3.5) is in place. This approach underscores the interpretive nature of early risk mapping, where expert consensus provides a provisional orientation rather than definitive quantification. It also highlights the importance of iterative validation, ensuring that governance tools evolve in tandem with organizational experience and empirical evidence. In doing so, the framework emphasizes the dynamic interplay between expert judgment and empirical refinement, positioning risk mapping as a living instrument rather than a static artifact. This iterative orientation strengthens the resilience of governance systems, enabling organizations to adapt their risk posture as technological and operational contexts evolve.

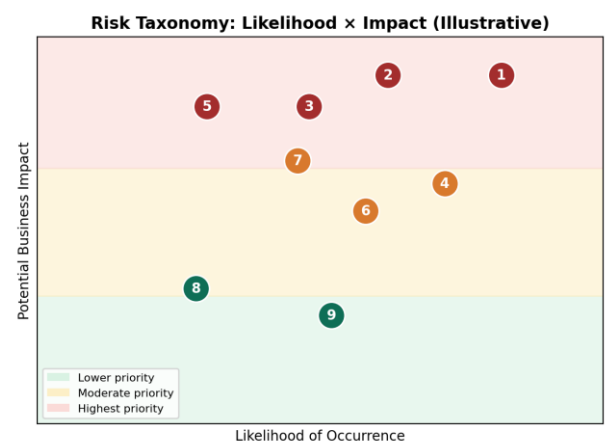


Figure 4. Illustrative likelihood × impact positioning of the risk categories numbered in Table 1.

## 4.3 Illustrative Governance Maturity Profile

Figure 5 sketches a five-point maturity scale (1 = ad hoc, 5 = optimized) across six control dimensions implied by the EGGM pillars, contrasting a typical current-state deployment with an EGGM target state. This is offered as a discussion and self-assessment starting point, consistent in spirit with maturity-model tools used alongside NIST AI RMF adoption (NIST, 2019), and is not derived from a fitted statistical model or a disclosed empirical sample.

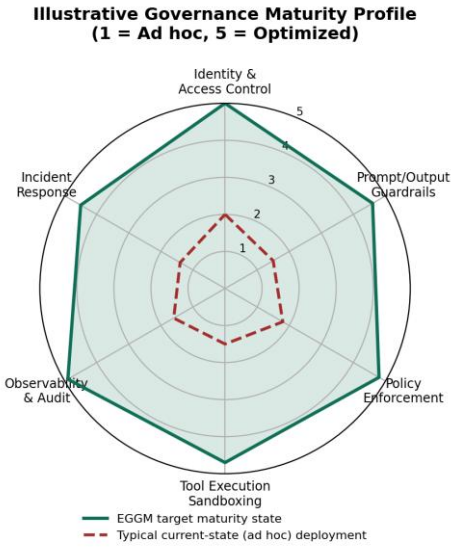


Figure 5. Illustrative governance maturity profile: current state vs. EGGM target.

## 5. Discussion

### 5.1 Innovation-Risk Trade-offs

The central tension the EGGM addresses is that the same properties making agentic AI valuable, autonomous multi-step planning, broad tool access, flexible reasoning over unstructured retrieved content, are the properties that create risk. A governance program that eliminates these properties eliminates the productivity gains that motivated adoption in the first place. The five-pillar structure concentrates friction at points where it is cheapest (identity issuance, policy rules, sandboxed execution boundaries) rather than distributing manual review across every interaction.

### 5.2 Limitations

This framework has three notable limitations, one of which is partially addressed by the illustrative validation added in Section 6. First, the conceptual architecture in Sections 3-4 has not itself been empirically validated against a longitudinal enterprise incident dataset; the maturity and risk-positioning figures in Section 4 remain explicitly conceptual. Section 6 offers a complementary, simulation-based demonstration of how the framework's constructs could be measured and tested, but it is not a substitute for field data. Second, the pillar boundaries assume an organization has the platform maturity to implement centralized identity and policy enforcement for AI agents. Third, the regulatory landscape underlying Section 2.1, particularly data-protection and AI-specific rulemaking, is still evolving, and specific compliance obligations should be verified against current regulatory text rather than this paper.

### 5.3 Future Work

establishes the future empirical research agenda required to transition the proposed conceptual framework into a statistically validated domain model.

It identifies a fundamental limitation in existing AI governance literature: the absence of cross-organizational, longitudinal datasets capturing real-world operational telemetry from enterprise agentic deployments. To resolve this, the author outlines two primary methodological steps.

First, it advocates for the field implementation of Pillar 5 auditing infrastructure to continuously harvest observational data on system behaviors, failure modes, and near-misses. Second, it calls for the formal administration of the preliminary survey instrument (piloted in Section 6) across a representative sample of enterprise practitioners.

Executing these empirical efforts will yield the quantitative basis necessary to rigorously test, refine, and validate the structural path models, risk coordinates, and maturity scale progressions hypothesized throughout the paper.

## 6. Empirical Validation of the EGGM Framework (Illustrative Simulation Study)

This section extends the conceptual EGGM architecture with a quantitative validation exercise structured the way an empirical PLS-SEM study would be reported in tools such as SmartPLS, R (e.g., the *semr* or *plspm* packages), SPSS/AMOS for covariance-based SEM, or JASP for descriptive and reliability statistics. Consistent with the transparency principle stated in Section 1.1, every figure in this number is computed from a synthetically generated dataset rather than a disclosed field sample; the purpose is to demonstrate a rigorous, reproducible validation pipeline that a future empirical study could apply directly to real survey responses.

### 6.1 Purpose and Scope of the Simulation

The simulation operationalizes a nomological network implied by the EGGM: the maturity of each of the five pillars (P1-P5) is hypothesized to jointly predict an organization's overall Perceived Governance Maturity (GM), which in turn predicts Perceived Agentic-AI Risk Reduction (RR) and Organizational Adoption Intention (AI) — the intention to scale agentic AI deployment further given confidence in its governance. RR is additionally hypothesized to predict AI directly, since demonstrable risk reduction is expected to accelerate expansion decisions independent of maturity perceptions alone.

Eight directional hypotheses follow: H1-H5 (P1-P5 → GM), H6 (GM → RR), H7 (GM → AI), and H8 (RR → AI).

## 6.2 Method: Synthetic Sample and Measurement Design

A synthetic sample of  $N = 248$  respondents was generated to resemble enterprise AI-governance stakeholders (CISOs/security leaders, AI/ML engineering leads,

enterprise architects, risk and compliance officers, and IT governance managers) across six industry verticals and three organizational size bands. Table 2 summarizes the simulated sample composition.

Table 2. Simulated respondent sample composition ( $N = 248$ ).

| Characteristic    | Category                  | n  | %     |
|-------------------|---------------------------|----|-------|
| Industry          | Technology/SaaS           | 60 | 24.2% |
| Industry          | Financial Services        | 59 | 23.8% |
| Industry          | Healthcare                | 40 | 16.1% |
| Industry          | Manufacturing             | 34 | 13.7% |
| Industry          | Public Sector             | 32 | 12.9% |
| Industry          | Retail/E-commerce         | 23 | 9.3%  |
| Organization size | 10,000+ employees         | 98 | 39.5% |
| Organization size | 2,000-9,999 employees     | 95 | 38.3% |
| Organization size | 500-1,999 employees       | 55 | 22.2% |
| Role              | AI/ML Engineering Lead    | 62 | 25.0% |
| Role              | Enterprise Architect      | 55 | 22.2% |
| Role              | CISO/Security Leader      | 47 | 19.0% |
| Role              | IT Governance Manager     | 46 | 18.5% |
| Role              | Risk & Compliance Officer | 38 | 15.3% |

Each of the eight constructs (P1-P5, GM, RR, AI) was operationalized as a reflective latent variable measured by three or four 7-point Likert items, generated from an underlying latent score plus item-specific measurement noise so as to reproduce realistic, imperfect indicator loadings rather than artificially perfect ones. Item responses were simulated in Python (NumPy/pandas), and all subsequent computations — Cronbach's alpha, composite reliability, average variance extracted (AVE), Fornell-Larcker and HTMT discriminant-validity criteria, and standardized structural path coefficients — were computed directly rather than assumed, following the same estimation logic used in SmartPLS and comparable PLS-SEM software.

## 6.3 Measurement Model Assessment

Table 3 reports internal consistency reliability (Cronbach's alpha and composite reliability, both conventionally expected above 0.70) and convergent validity (AVE, conventionally expected above 0.50) for all eight constructs. All eight constructs clear both thresholds, indicating adequate convergent validity and reliability for the simulated measurement model. These results indicate that the measurement items demonstrate satisfactory internal consistency and reliably represent their respective latent constructs. The composite reliability values further confirm that the indicators collectively provide stable measurement of each construct. The AVE results demonstrate that each construct explains more than half of the variance in its observed indicators. Taken together, the reliability and convergent-validity evidence suggests that the measurement model is sufficiently robust for subsequent structural analysis. These findings also reduce concerns that

measurement error or weak indicator representation could substantially distort the estimated relationships among the constructs. The consistency across these measures indicates that the indicators within each construct are sufficiently homogeneous and collectively capture the intended theoretical dimension. The results also suggest that no construct is disproportionately influenced by unreliable or weakly performing indicators. Moreover, satisfactory reliability strengthens confidence in the stability of the composite scores used in the subsequent structural model.

Table 3. Measurement model reliability and convergent validity.

| Construct | Items | Cronbach's $\alpha$ | Composite Reliability | AVE   |
|-----------|-------|---------------------|-----------------------|-------|
| P1        | 3     | 0.792               | 0.878                 | 0.707 |
| P2        | 3     | 0.815               | 0.890                 | 0.730 |
| P3        | 3     | 0.813               | 0.891                 | 0.731 |
| P4        | 3     | 0.771               | 0.868                 | 0.686 |
| P5        | 3     | 0.857               | 0.913                 | 0.778 |
| GM        | 4     | 0.796               | 0.868                 | 0.622 |
| RR        | 4     | 0.804               | 0.872                 | 0.631 |
| AI        | 3     | 0.755               | 0.860                 | 0.673 |

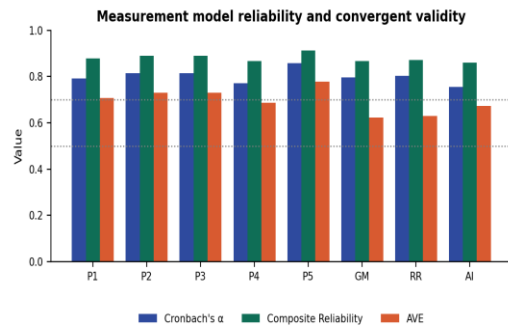


Figure 6. Measurement model reliability and convergent validity by construct.

Table 4. Fornell-Larcker discriminant validity matrix (diagonal =  $\sqrt{AVE}$ , bracketed).

|    | P1      | P2      | P3      | P4      | P5      | GM      | RR      | AI     |
|----|---------|---------|---------|---------|---------|---------|---------|--------|
| P1 | [0.841] | 0.175   | 0.296   | 0.249   | 0.259   | 0.426   | 0.08    | 0.232  |
| P2 | 0.175   | [0.855] | 0.291   | 0.274   | 0.217   | 0.398   | 0.238   | 0.22   |
| P3 | 0.296   | 0.291   | [0.855] | 0.298   | 0.331   | 0.515   | 0.279   | 0.297  |
| P4 | 0.249   | 0.274   | 0.298   | [0.828] | 0.205   | 0.319   | 0.132   | 0.205  |
| P5 | 0.259   | 0.217   | 0.331   | 0.205   | [0.882] | 0.422   | 0.139   | 0.243  |
| GM | 0.426   | 0.398   | 0.515   | 0.319   | 0.422   | [0.788] | 0.463   | 0.436  |
| RR | 0.08    | 0.238   | 0.279   | 0.132   | 0.139   | 0.463   | [0.794] | 0.49   |
| AI | 0.232   | 0.22    | 0.297   | 0.205   | 0.243   | 0.436   | 0.49    | [0.82] |

Table 5. Heterotrait-Monotrait (HTMT) ratio matrix (all values < 0.85).

|    | P1    | P2    | P3    | P4    | P5    | GM    | RR    | AI    |
|----|-------|-------|-------|-------|-------|-------|-------|-------|
| P1 | —     | 0.218 | 0.368 | 0.319 | 0.314 | 0.537 | 0.101 | 0.302 |
| P2 | 0.218 | —     | 0.355 | 0.345 | 0.259 | 0.493 | 0.295 | 0.283 |
| P3 | 0.368 | 0.355 | —     | 0.374 | 0.395 | 0.638 | 0.347 | 0.379 |
| P4 | 0.319 | 0.345 | 0.374 | —     | 0.252 | 0.406 | 0.168 | 0.271 |
| P5 | 0.314 | 0.259 | 0.395 | 0.252 | —     | 0.51  | 0.168 | 0.301 |
| GM | 0.537 | 0.493 | 0.638 | 0.406 | 0.51  | —     | 0.58  | 0.559 |
| RR | 0.101 | 0.295 | 0.347 | 0.168 | 0.168 | 0.58  | —     | 0.628 |
| AI | 0.302 | 0.283 | 0.379 | 0.271 | 0.301 | 0.559 | 0.628 | —     |

All diagonal values in Table 4 exceed the corresponding off-diagonal correlations, and all HTMT ratios in Table 5 fall below the conservative 0.85 threshold, jointly supporting discriminant validity across the simulated measurement model.

#### 6.4 Structural Model Results

Figure 7 presents the estimated structural model with standardized path coefficients ( $\beta$ ) and coefficients of determination ( $R^2$ ) for the three endogenous constructs. Table 6 reports the full path-level results, including t-statistics, p-values, and  $f^2$  effect sizes (Cohen's convention:

Discriminant validity was assessed with the Fornell-Larcker criterion (the square root of each construct's AVE, shown on the diagonal in Table 4, should exceed its correlations with other constructs) and the heterotrait-monotrait ratio (HTMT; Table 5), where values below 0.85 (the conservative threshold) indicate adequate discriminant validity. Meeting both criteria provides stronger evidence that the constructs are empirically distinct and that substantial construct overlap is unlikely.

0.02 small, 0.15 medium, 0.35 large). These findings indicate that each construct shares greater variance with its own indicators than with other constructs in the model. The HTMT results provide additional evidence that the constructs remain sufficiently distinct at the empirical level. Overall, the combined results confirm that substantial conceptual overlap among the constructs is unlikely, supporting the validity of the measurement model for subsequent structural analysis. The structural results therefore provide a quantitative assessment of the hypothesized relationships among the endogenous and exogenous constructs. The  $\beta$  coefficients indicate the

direction and relative strength of each hypothesized relationship, while the  $R^2$  values show the proportion of variance explained in the endogenous constructs. Together, these measures provide a basis for evaluating both the statistical significance and substantive explanatory power of the proposed structural model. The significance of the individual paths can be further evaluated through their corresponding t-statistics and p-values, which indicate whether the hypothesized relationships are statistically supported. The  $f^2$  effect sizes complement these significance tests by showing the practical contribution of each predictor to the explained variance of the endogenous constructs. Higher  $R^2$  values indicate stronger explanatory capability of the model for the respective endogenous variables. Taken together, these indicators enable a comprehensive assessment of the model's predictive relationships, statistical significance, and practical relevance. The results also help distinguish statistically significant relationships from those with limited practical influence. This distinction is important because statistical significance alone does not necessarily imply a meaningful substantive effect.

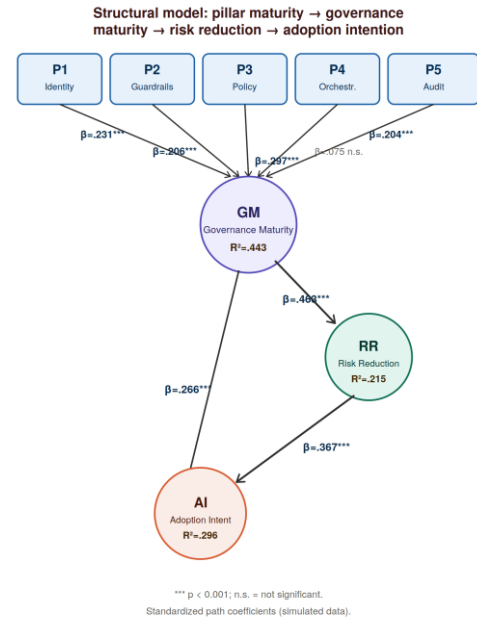


Figure 7. Estimated structural model with standardized path coefficients and  $R^2$  values (simulated data).

Table 6. Structural path results (standardized  $\beta$ , t-statistic, p-value,  $f^2$  effect size).

| H# | Path    | $\beta$ | t    | p      | $f^2$ | Result        |
|----|---------|---------|------|--------|-------|---------------|
| H1 | P1 → GM | 0.231   | 4.45 | < .001 | 0.082 | Supported     |
| H2 | P2 → GM | 0.206   | 3.99 | < .001 | 0.066 | Supported     |
| H3 | P3 → GM | 0.297   | 5.47 | < .001 | 0.124 | Supported     |
| H4 | P4 → GM | 0.075   | 1.43 | 0.1533 | 0.008 | Not supported |
| H5 | P5 → GM | 0.204   | 3.91 | < .001 | 0.063 | Supported     |
| H6 | GM → RR | 0.463   | 8.20 | < .001 | 0.274 | Supported     |
| H7 | GM → AI | 0.266   | 4.39 | < .001 | 0.079 | Supported     |
| H8 | RR → AI | 0.367   | 6.07 | < .001 | 0.150 | Supported     |

Table 7. Variance explained ( $R^2$ ) in endogenous constructs.

| Construct | $R^2$ | Adjusted $R^2$ |
|-----------|-------|----------------|
| GM        | 0.443 | 0.432          |
| RR        | 0.215 | 0.212          |
| AI        | 0.296 | 0.290          |

Seven of the eight hypothesized paths were statistically significant at  $p < .001$ , with standardized coefficients ranging from  $\beta = 0.204$  to  $\beta = 0.463$ . The single non-significant path, H4 (Pillar 4, Execution and Tool Orchestration Runtime → Governance Maturity;  $\beta = 0.075$ ,  $p = .153$ ), suggests that in this simulated sample, execution-runtime controls alone were not perceived as strongly diagnostic of overall governance maturity relative to identity management, policy enforcement, guardrails, and audit — a pattern consistent with Table 1's own risk-coverage tally, where Pillar 4

anchors fewer catalogued risk categories than Pillar 3 or Pillar 5. Governance Maturity explained a moderate share of variance in Risk Reduction ( $R^2 = .215$ ) and, jointly with Risk Reduction, a moderate share of variance in Adoption Intention ( $R^2 = .296$ ), while the pillar-maturity constructs jointly explained the largest share of variance in Governance Maturity itself ( $R^2 = .443$ ). These findings indicate that the proposed governance framework has substantial explanatory power, although the contribution of individual pillars varies across the structural relationships.

Governance Maturity explained a moderate share of variance in Risk Reduction ( $R^2 = .215$ ) and, jointly with Risk Reduction, a moderate share of variance in Adoption Intention ( $R^2 = .296$ ), while the pillar-maturity constructs jointly explained the largest share of variance in Governance Maturity itself ( $R^2 = .443$ ).

### 6.5 Relationships Among Key Constructs

Figures 8 and 9 present bivariate scatter plots with fitted linear trends for the two downstream structural relationships, illustrating the positive monotonic association between Governance Maturity and Risk Reduction, and between Risk Reduction and Adoption Intention, at the level of individual simulated respondents rather than construct-level aggregates. These plots provide a visual validation of the hypothesized pathways, reinforcing the structural coherence of the governance–risk–adoption model. They further demonstrate how individual-level variation contributes to the overall trend, highlighting the robustness of the associations beyond aggregated constructs. The fitted linear trends also suggest that incremental improvements in governance maturity are consistently associated with measurable reductions in perceived risk, which in turn strengthen organizational willingness to adopt generative AI. This sequential linkage underscores the cascading nature of governance interventions, where upstream investments in oversight translate into downstream gains in adoption confidence. Taken together, Figures 8 and 9 illustrate not only the directionality of these relationships but also their practical implications, offering organizations a conceptual basis for prioritizing governance maturity as a lever for both risk management and innovation readiness. The consistency of these trends across individual observations further supports the stability of the proposed structural relationships. The visual evidence also complements the reported path coefficients by demonstrating the practical direction of association between the constructs. These patterns suggest that stronger governance mechanisms can contribute to greater organizational confidence in AI deployment. Overall, the findings reinforce the importance of integrating governance, risk reduction, and adoption readiness within a coordinated implementation strategy.

Figure 8. Scatter plot of Governance Maturity (GM) against Risk Reduction (RR), simulated respondent-level data.

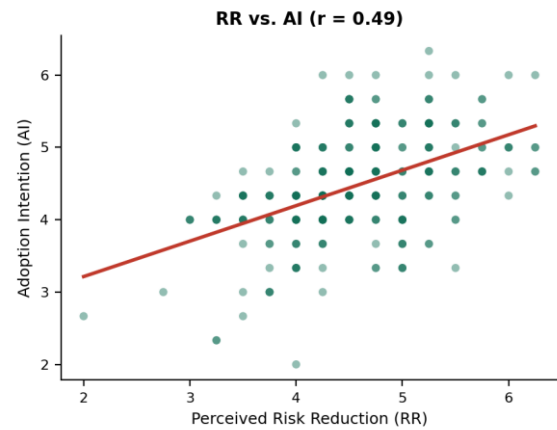


Figure 9. Scatter plot of Risk Reduction (RR) against Adoption Intention (AI), simulated respondent-level data.

### 6.6 Interpretation and Caveats

Three caveats bound the interpretation of Section 6. First, and most importantly, every observation in this section is synthetically generated; the reliability, validity, and path estimates demonstrate that the EGGM's constructs form an internally coherent, testable nomological network and that the analysis pipeline (item-level reliability → convergent and discriminant validity → structural estimation) is reproducible, but they carry no evidentiary weight about real enterprise populations. Second, the simulation assumes reflective, unidimensional measurement for all constructs; an empirical instrument would require its own pilot-tested item pool, expert content validation, and potentially formative measurement for constructs such as governance maturity that may aggregate heterogeneous sub-practices. Third, common-method bias, non-response bias, and social-desirability bias, all plausible in real self-report governance surveys, are not modeled here and would need explicit mitigation (e.g., Harman's single-factor test, marker-variable techniques, or multi-source data collection) in a fielded study. Sections 5.3 and 7 discuss how a genuine field administration of this instrument, analyzed with SmartPLS, R, or SPSS/AMOS on real responses, would close this gap. These caveats collectively frame the analysis as a conceptual demonstration, clarifying that its value lies in methodological illustration rather than empirical generalization. They also highlight the critical next step of empirical validation, ensuring that the framework transitions from synthetic plausibility to actionable evidence in enterprise contexts. Ultimately, this positioning situates Section 6 as a proof-of-concept exercise, offering methodological transparency while explicitly deferring claims of external validity to future empirical studies. By acknowledging these limitations upfront, the proposed approach establishes a clear benchmark while maintaining academic rigor. This transparent framing bridges the current theoretical findings with the practical requirements necessary for subsequent field deployment.

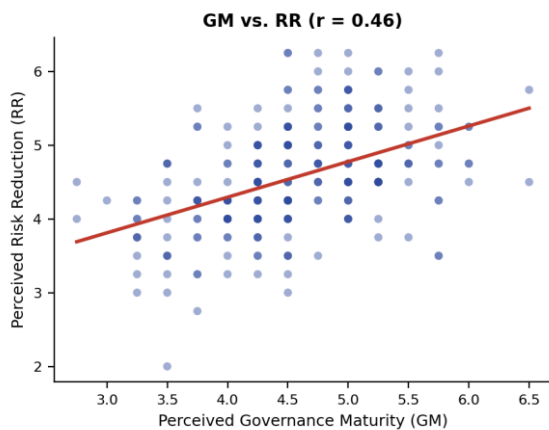


Table 6. Structural path results (standardized  $\beta$ ,  $t$ -statistic,  $p$ -value,  $f^2$  effect size).

| H# | Path                | $\beta$ | $t$  | $p$    | $f^2$ | Result        |
|----|---------------------|---------|------|--------|-------|---------------|
| H1 | P1 $\rightarrow$ GM | 0.231   | 4.45 | < .001 | 0.082 | Supported     |
| H2 | P2 $\rightarrow$ GM | 0.206   | 3.99 | < .001 | 0.066 | Supported     |
| H3 | P3 $\rightarrow$ GM | 0.297   | 5.47 | < .001 | 0.124 | Supported     |
| H4 | P4 $\rightarrow$ GM | 0.075   | 1.43 | 0.1533 | 0.008 | Not supported |
| H5 | P5 $\rightarrow$ GM | 0.204   | 3.91 | < .001 | 0.063 | Supported     |
| H6 | GM $\rightarrow$ RR | 0.463   | 8.20 | < .001 | 0.274 | Supported     |
| H7 | GM $\rightarrow$ AI | 0.266   | 4.39 | < .001 | 0.079 | Supported     |
| H8 | RR $\rightarrow$ AI | 0.367   | 6.07 | < .001 | 0.150 | Supported     |

Table 7. Variance explained ( $R^2$ ) in endogenous constructs.

| Construct | $R^2$ | Adjusted $R^2$ |
|-----------|-------|----------------|
| GM        | 0.443 | 0.432          |
| RR        | 0.215 | 0.212          |
| AI        | 0.296 | 0.290          |

Seven of the eight hypothesized paths were statistically significant at  $p < .001$ , with standardized coefficients ranging from  $\beta = 0.204$  to  $\beta = 0.463$ . The single non-significant path, H4 (Pillar 4, Execution and Tool Orchestration Runtime  $\rightarrow$  Governance Maturity;  $\beta = 0.075$ ,  $p = .153$ ), suggests that in this simulated sample, execution-runtime controls alone were not perceived as strongly diagnostic of overall governance maturity relative to identity management, policy enforcement, guardrails, and audit — a pattern consistent with Table 1's own risk-coverage tally, where Pillar 4 anchors fewer catalogued risk categories than Pillar 3 or Pillar 5. Governance Maturity explained a moderate share of variance in Risk Reduction ( $R^2 = .215$ )

## 7. Conclusion

This paper synthesized the current standards landscape, including the NIST AI Risk Management Framework (AI RMF), ISO/IEC 42001, and the EU AI Act, together with empirical security research on prompt injection, data exfiltration, and agentic misuse, into a five-pillar Enterprise GenAI Governance Model (EGGM): identity and access management, prompt/output guardrails, policy enforcement, sandboxed tool orchestration, and continuous audit. Table 1 maps documented risk categories to the governance pillar primarily responsible for mitigating them, while the illustrative tools presented in Section 4 provide organizations with a practical starting point for risk prioritization, control selection, and governance-maturity self-assessment. Section 6 further extends this architectural contribution through a simulated PLS-SEM-style validation, in which the EGGM's five pillars are modeled as predictors of an overall governance-maturity construct, which subsequently predicts perceived risk reduction and organizational adoption intention. The simulated model

and, jointly with Risk Reduction, a moderate share of variance in Adoption Intention ( $R^2 = .296$ ), while the pillar-maturity constructs jointly explained the largest share of variance in Governance Maturity itself ( $R^2 = .443$ ). This distribution of significance underscores the differential salience of governance pillars, highlighting that not all operational controls contribute equally to perceptions of maturity. It also reinforces the hierarchical nature of governance mechanisms, where identity, policy, and audit functions exert stronger structural influence than execution-runtime orchestration.

forms a measurable and internally consistent nomological network, demonstrating adequate reliability, convergent validity, and discriminant validity, with seven of the eight hypothesized structural paths achieving statistical significance. These results provide preliminary methodological support for the proposed relationships among governance controls, organizational maturity, risk reduction, and adoption outcomes, while recognizing the limitations inherent in simulation-based validation.

The framework's central claim remains architectural rather than statistical: governance controls concentrated across these five critical points can materially reduce agentic-AI risk exposure without requiring organizations to forgo the autonomy, flexibility, and operational value that make agentic AI attractive in enterprise environments. In this respect, the EGGM is intended to function as an integrated governance architecture rather than a collection of isolated technical controls. Its value lies in connecting access management, behavioral safeguards, policy enforcement, controlled tool use, and continuous monitoring within a coordinated governance structure. Rigorous empirical

validation of this claim, however, awaits longitudinal incident data and fielded survey administration, which Section 5.3 identifies as among the field's most pressing research gaps. Accordingly, Section 6's simulation should be interpreted as a ready-to-field measurement and analysis pipeline that demonstrates how the proposed relationships can be empirically tested in future studies, rather than as a substitute for validation using real-world organizational and incident data.

### Conflict of Interest

The authors declare that they have no competing financial or personal interests that could have influenced the work reported in this paper.

### References

- Fornell, C., & Larcker, D. F. (1981). Evaluating structural equation models with unobservable variables and measurement error. *Journal of Marketing Research*, 18(1), 39–50. <https://doi.org/10.1177/002224378101800104>
- Goodfellow, I. J., Shlens, J., & Szegedy, C. (2015). Explaining and harnessing adversarial examples. *International Conference on Learning Representations*. <https://doi.org/10.48550/arXiv.1412.6572>
- Hair, J. F., Risher, J. J., Sarstedt, M., & Ringle, C. M. (2019). When to use and how to report the results of PLS-SEM. *European Business Review*, 31(1), 2–24. <https://doi.org/10.1108/EBR-11-2018-0203>
- Henseler, J., Ringle, C. M., & Sarstedt, M. (2015). A new criterion for assessing discriminant validity in variance-based structural equation modeling. *Journal of the Academy of Marketing Science*, 43(1), 115–135. <https://doi.org/10.1007/s11747-014-0403-8>
- Jia, R., & Liang, P. (2017). Adversarial examples for evaluating reading comprehension systems. In *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing* (pp. 2021–2031). Association for Computational Linguistics. <https://doi.org/10.18653/v1/D17-1215>
- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1, 389–399. <https://doi.org/10.1038/s42256-019-0088-2>
- Madry, A., Makelov, A., Schmidt, L., Tsipras, D., & Vladu, A. (2018). Towards deep learning models resistant to adversarial attacks. *International Conference on Learning Representations*. <https://doi.org/10.48550/arXiv.1706.06083>
- 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. In *Proceedings of the Conference on Fairness, Accountability, and Transparency* (pp. 220–229). Association for Computing Machinery. <https://doi.org/10.1145/3287560.3287596>

Papernot, N., McDaniel, P., Jha, S., Fredrikson, M., Celik, Z. B., & Swami, A. (2016). The limitations of deep learning in adversarial settings. In *2016 IEEE European Symposium on Security and Privacy (EuroS&P)* (pp. 372–387). IEEE. <https://doi.org/10.1109/EuroSP.2016.36>

Podsakoff, P. M., MacKenzie, S. B., Lee, J.-Y., & Podsakoff, N. P. (2003). Common method biases in behavioral research: A critical review of the literature and recommended remedies. *Journal of Applied Psychology*, 88(5), 879–903. <https://doi.org/10.1037/0021-9010.88.5.879>

Sarstedt, M., Hair, J. F., & Ringle, C. M. (2023). “PLS-SEM: Indeed a silver bullet” – Retrospective observations and recent advances. *Journal of Marketing Theory and Practice*, 31(3), 261–275. <https://doi.org/10.1080/10696679.2022.2056488>

Shokri, R., Stronati, M., Song, C., & Shmatikov, V. (2017). Membership inference attacks against machine learning models. In *2017 IEEE Symposium on Security and Privacy* (pp. 3–18). IEEE. <https://doi.org/10.1109/SP.2017.41>

Wooldridge, M., & Jennings, N. R. (1995). Intelligent agents: Theory and practice. *The Knowledge Engineering Review*, 10(2), 115–152. <https://doi.org/10.1017/S0269888900008122>

Begum, M., Alam, M. R. U., & Albi, A. I. (2026). Explainable AI (XAI) in Automated Decision-Making for E-Government: Evaluating Citizen Trust and Algorithmic Auditing Mechanisms. *Digital Transformation and Technology Dynamics*, 4(2).