

COMPREHENSIVE REPORT & LEGAL RISK ASSESSMENT

AI Governance in Regulated Industries

Closing the Lifecycle Accountability Gap in Financial Services and Healthcare

Joshua M. Reh

Principal — AI Advisory, Cijara Group

cijaragroup.com · joshua@cijaragroup.com

Prepared for Boards, Chief Risk Officers, and General Counsel

Synthesizing five source documents and a curated source library
on enterprise AI governance, agentic autonomy, and regulation

July 31, 2026 · Version 1

Contents

1. Executive Summary
2. The Regulatory Landscape: Fragmentation as the Defining Condition
 - 2.1 Global divergence: EU, United States, China
 - 2.2 The shift to lifecycle governance
 - 2.3 Sector-specific high-risk classifications under the EU AI Act
3. Financial Services: From Periodic Validation to Continuous Oversight
4. Healthcare: Fragmented Regulation and the Reinsertion of Human Judgment
5. Governance Strategies That Are Working
6. Technical Architecture for Agentic Governance
7. Legal Risk Assessment
 - 7.1 Methodology
 - 7.2 Risk register and scoring
 - 7.3 Detailed assessments of the top risks
 - 7.4 Escalation and outside-counsel criteria
8. Implementation Roadmap
9. Conclusion
- References
- Appendix: Curated Source Library

1. Executive Summary

The rapid deployment of artificial intelligence — particularly *agentic* systems that act autonomously across enterprise software — has opened a structural **accountability gap** in the financial services and healthcare sectors. Agentic systems execute thousands of autonomous decisions daily, instantly breaking legacy governance models built for quarterly review cycles^[5]. Institutional adoption continues to accelerate, yet the governance frameworks overseeing these systems remain tethered to point-in-time compliance models: periodic model validation, static checklists, and committee-driven review cycles designed for the Basel II era^{[1][2][5]}. This report synthesizes five source documents — a cross-sector governance analysis, a regulated-industry research brief, the *2026 Blueprint for Enterprise AI Governance*, *Governing the Agentic Enterprise*, and *The Accountability Gap: How Highly Regulated Industries are Governing Agentic AI* — together with a curated source library of regulations, standards, academic research, and industry guidance, and applies a formal severity-by-likelihood legal risk assessment to the exposures identified.

Key findings

- **Deployment is outpacing accountability.** Only about one-third of organizations report governance maturity at level three or higher for agentic AI controls, despite widespread integration into high-stakes workflows^[1]. The 2026 Blueprint quantifies the same "readiness chasm": 70% of enterprises have established cross-functional AI oversight committees on paper, but only 14% report full preparedness to deploy AI capabilities, and 31% say they are struggling to keep pace with preparation requirements^[3].
- **Sector convergence on lifecycle governance.** Finance and healthcare are converging on continuous, post-deployment monitoring — but from opposite directions. Finance is adapting 15-year-old model risk scaffolds (notably SR 11-7), while healthcare is improvising by legislating human judgment back into automated processes^[1].
- **Regulatory fragmentation is the operating condition.** In the absence of U.S. federal AI legislation, state-level activity is surging — more than 250 health AI bills were introduced in 2025 alone — while the EU AI Act, Singapore MAS guidelines, Vietnam's Law on AI, and the ASEAN framework create a "contested perimeter" for multi-jurisdiction institutions^{[1][3]}.
- **The risk profile has changed in kind, not degree.** Traditional generative AI risks (hallucination, bias, misinformation) are being displaced by agentic risks: irreversible actions, state changes, and corrupted production data. Controls must shift accordingly, from output filtering to action bounding^{[3][4]}.
- **Trust has measurable ROI.** Organizations prioritizing trustworthy AI and human oversight are 60% more likely to double their return on AI investment than those focused solely on cost-cutting or automation^[1].

- **Governance gaps are a leading cause of deployment failure.** Gartner predicts governance gaps will cause 50% of AI agent deployment failures by 2030^[5].
- **The dashboard illusion.** Most organizations rely on governance platforms that stop at the dashboard — but static dashboards cannot enforce runtime policies or stop an autonomous agent mid-execution^[5].
- **Liability transfers upward, never away.** Regulators consistently treat the organization as responsible for an agent's actions exactly as if an employee had taken them; if an agent calls a tool, moves data, or initiates a transaction, the institution holds ultimate accountability for that execution^[5].

Headline risk conclusions

Applying the severity × likelihood framework detailed in Section 7, three risk clusters score in the **RED (critical)** band and demand immediate executive attention: algorithmic claim denials in payer healthcare (score 25), uncontrolled agentic model drift under eroding SR 11-7 assumptions (score 16), shadow AI and rogue-agent exposure (score 16), and the documentation/audit-readiness gap (score 16). A further cluster — including unbounded autonomous actions and regulatory fragmentation — sits at the top of the **ORANGE (high)** band. None of these risks is adequately addressed by legacy, point-in-time governance; all of them are substantially mitigated by the continuous-runtime governance, risk-tiered autonomy limits, and infrastructure-level controls described in Sections 5 and 6.

The strategic thesis in one sentence. The organizations treating AI governance as organizational design will set the template; those treating it as point-in-time legal compliance will become the case studies^[1].

2. The Regulatory Landscape: Fragmentation as the Defining Condition

2.1 Global divergence: EU, United States, China

Enterprises must navigate a fragmented and competing global regulatory landscape. The 2026 Blueprint characterizes three regulatory archetypes^[3]:

Table 1 Global AI regulatory archetypes (synthesized from the 2026 Blueprint)

	European Union — "The Standard-Setter"	United States — "The Agile Innovator"	China — "The State Controller"
Core philosophy	Risk-based safeguards for fundamental rights	Innovation-first; voluntary compliance and agency directives	Balancing technological dominance with political alignment

Regulatory style	Comprehensive, centralized Act categorizing societal risk	Piecemeal, sector-specific bills and Executive Orders	Agile, targeted policies adapting swiftly to new technology (e.g., generative AI)
Impact	Strict compliance burden; outright bans on high-risk use cases (e.g., social scoring)	High flexibility, but risks losing global influence without a centralized framework	Favorable to adoption and copyright, but strict censorship and state-control requirements remain

For U.S.-centric institutions, the domestic picture is complicated by state-level fragmentation: with no federal AI statute, states are legislating independently, producing divergent obligations on disclosure, payer accountability, and therapeutic AI that a multi-state institution must satisfy simultaneously^[1].

2.2 The shift to lifecycle governance

Across jurisdictions, regulators are converging on one substantive expectation: **lifecycle governance** — continuous post-deployment monitoring rather than one-time approval. Key milestones identified in the source material^[1]:

- **European Union:** the EU AI Act (high-risk obligations enforceable from August 2026) applies regardless of where an institution is established if it interacts with EU users.
- **Singapore (MAS):** proposed guidelines (November 2025) require board-level accountability and cross-functional oversight committees.
- **Southeast Asia:** Vietnam's Law on AI (effective March 2026) and the upcoming ASEAN Digital Economy Framework Agreement (late 2026) are creating binding interoperability rules.
- **United States (healthcare):** a deregulatory federal posture on tools — the FDA's revised Clinical Decision Support guidance of January 6, 2026 expands the software categories reaching market without premarket review — coexists with adversarial federalism toward state structures governing consequences, and with documented enforcement-capacity strain at the FDA following 2025 staffing cuts^[1].

Beyond statutes, a layer of standards and frameworks is hardening into the practical definition of compliance. The curated source library assembled for this report identifies the instruments regulators and auditors actually reference^[6]: the EU AI Act's Article 14 human-oversight requirement; the NIST AI Risk Management Framework (AI RMF 1.0) and its forthcoming Trustworthy AI in Critical Infrastructure profile; ISO/IEC 42001 for AI Management Systems (AIMS); Singapore IMDA's Model AI Governance Framework for Agentic AI and the MAS safeguards for agentic finance at runtime; Japan's second AI Basic Plan; and the Hiroshima AI Process reporting framework. All pathways point to the same destination: **continuous lifecycle accountability**, with oversight controls, confidence thresholds, and upfront risk-bounding as common requirements^[5].

2.3 Sector-specific high-risk classifications under the EU AI Act

The EU AI Act's high-risk regime lands differently across sectors — a critical input to any institution's exposure mapping^[5]:

Table 2 Sector-specific high-risk classifications (EU AI Act)

Sector	High-risk classification and obligations
Financial services	Credit scoring and lending AI are classified as high-risk; requires strict algorithmic-bias scrutiny and integration with MiFID II and AML directives
Human resources & employment	Recruitment and performance-evaluation systems require fundamental-rights safeguards and explainability
Healthcare & medical devices	Dual compliance required (Medical Device Regulation and EU AI Act); impacts clinical decision support and diagnostic AI
Law enforcement	Judicial decision-support requires extensive transparency and human oversight

Implication for legal teams. Because the rules are fragmented and shifting, governance cannot be built around any single rule set. Institutions must build the reporting lines, monitoring infrastructure, and cross-functional authority required to answer *any supervisor in any jurisdiction at any time*^[1]. Regulators are writing rules for 2026 while most institutions run governance on a 2018 operating model — the organizations that will comfortably pass 2027 conformity assessments are those treating today's regulatory deferrals as build time^[5].

3. Financial Services: From Periodic Validation to Continuous Oversight

The financial sector faces a significant organizational gap: regulatory expectations are moving toward real-time surveillance, while internal governance functions still operate on quarterly or semi-annual review cycles^[1].

3.1 The erosion of SR 11-7

For nearly 15 years, U.S. banking model risk management has relied on SR 11-7, a framework premised on models that are discrete, stable, and periodically validated. That scaffold is strained by agentic AI — systems that autonomously parse trade data, resolve compliance exceptions, and initiate multi-step workflows. Because these systems recalibrate autonomously, material behavioral changes can occur *between* formal review events, defeating the periodic-validation assumption at the heart of the framework^[1]. The Blueprint's

mapping of KPMG's five governance pillars onto the NIST AI Risk Management Framework (Govern–Map–Measure–Manage) is one practical route for translating board-level strategy into operational risk machinery to fill this gap^[3].

3.2 Deployment trends and the risk-awareness gap

- **Frontier deployments are already live.** Goldman Sachs has embedded Anthropic engineers to co-develop autonomous agents for trade accounting and client vetting (KYC); these agents make "micro decisions" based on reasoning rather than static rules^[1].
- **The economic stakes are large.** BCG and OpenAI project that agentic AI could increase bank profitability by roughly 30% by 2030, with potential cost reductions of 30–40% through agentic AI integration^[1].
- **Awareness is not mitigation.** According to the 2025 Stanford AI Index survey cited in the source material, the widest gap in AI risk management — 15 percentage points — lies between considering "autonomous/unintended actions" relevant and actively mitigating them^[1].

3.3 Sector-specific strategies

The source material recommends three implementation priorities for financial institutions^[1]:

- **Living AI inventory:** automated scanning to detect "shadow deployments" across API endpoints and vendor integrations — you cannot govern what you cannot observe^[4].
- **Tiered validation:** calibrate validation frequency to a system's autonomy level; high-autonomy agents warrant shorter cycles or continuous "champion-challenger" testing.
- **Safety controls:** global kill switches and infrastructure-level circuit breakers capable of halting agents in seconds during an anomaly.

4. Healthcare: Fragmented Regulation and the Reinsertion of Human Judgment

The healthcare AI environment is defined by a federal posture simultaneously deregulatory toward tools and adversarial toward the state-level structures attempting to govern their consequences^[1].

4.1 The FDA's two-track posture

On January 6, 2026, the FDA issued revised guidance on Clinical Decision Support (CDS) software, expanding the category of software that can reach the market without premarket review — shifting the burden of validation onto health systems at the point of procurement. Conversely, the FDA is developing a

Total Product Lifecycle framework for high-risk AI medical devices, though 2025 staffing cuts have hindered the agency's enforcement and review capacity^[1].

4.2 The crisis in payer AI

The "payer lane" presents the most thoroughly documented evidence of harm in any sector covered by the source material^[1]:

- AI systems used in insurance utilization review for Medicare Advantage have an **82% overturn rate on appeal** — yet **fewer than 1% of patients contest denials**, so the system receives almost no corrective signal and care goes undelivered on the basis of incorrect algorithmic determinations.
- Three in four health plans use AI for prior authorization decisions.

4.3 The state-level response: legislating humans back in

Table 3 Four state legislative lanes in health AI (as catalogued in the source material)

Lane	Example	Substance
Mental health chatbots	California; Illinois	Prohibit AI from making independent therapeutic decisions without licensed professional review
Patient disclosure	Texas	Written disclosure to patients before AI is used in their treatment
Payer accountability	California SB 1120	Prohibits health plans from basing medical-necessity determinations solely on AI
Burden-shifting	Louisiana SB 246 (proposed)	Presumes any AI-influenced denial invalid unless the insurer proves independent human clinical judgment was used

4.4 Sector-specific strategies

- **Point-of-care disclosure:** build multi-state disclosure infrastructure into EHR workflows to satisfy varying state mandates (e.g., Texas and California) simultaneously^[1].
- **Physician-level review:** redesign utilization review so every adverse determination is signed off by a qualified human, anticipating burden-shifting statutes like Louisiana's^[1].
- **Vendor due diligence:** factor the FDA's capacity gap into procurement; do not rely on FDA clearance as a proxy for clinical safety^[1].

5. Governance Strategies That Are Working

The most effective organizations are shifting from governance as a point-in-time legal compliance exercise to governance as a continuous, operational, organizational design discipline^[2]. Four mutually reinforcing strategies dominate the evidence base.

5.1 Continuous "runtime" governance embedded in the platform

Instead of relying on written policy documents or quarterly review meetings, successful organizations embed governance directly into the AI platform's control plane^[2]:

- **Action-level enforcement:** every consequential action an agent attempts is checked against policy before execution; the system can automatically block, allow, or route the action to a human.
- **Shadow-mode deployments:** governance layers first run in "shadow mode" to observe actual AI behavior and how policies would have intercepted actions, allowing controls to be tuned on real behavior before live enforcement.
- **Continuous monitoring:** statistical (drift detection on input/output distributions), behavioral (telemetry on every agent action, tool call, and decision rationale in tamper-evident logs), and fairness monitoring (ongoing disparate-impact testing, especially for credit and coverage decisions) are built into the production stack^{[1][2]}.

5.2 Dynamic, "human-in-the-lead" oversight

Regulators demand meaningful human oversight, but reviewing every AI action creates alert fatigue and automation complacency. Leading organizations adopt strategic oversight models^[2]:

- **Risk-based routing:** confidence bands (green/amber/red lanes) route decisions — low-risk, high-confidence actions auto-approve; high-risk, irreversible decisions escalate to human queues under strict SLAs.
- **AI-instigated escalation:** advanced frameworks let the AI self-assess uncertainty, contextual failures, or ethical conflicts and autonomously trigger human intervention precisely when risk is highest. In practice this is implemented as **confidence-based escalation:** the AI operates autonomously until a low-confidence score or high-risk parameter triggers automatic escalation to human judgment, while human leaders monitor data pipelines, set the confidence thresholds, and manage systemic drift before algorithmic decay impacts the business^[5]. Academic work in the source library — on decoupled human-in-the-loop systems for controlled autonomy, AI-instigated human oversight in clinical AI, and a constitutive-vs-corrective taxonomy of human runtime involvement — is formalizing these patterns^[6].

- **Simulator training:** oversight is treated as an operational skill; reviewers train in simulator environments — much like aviation pilots — to practice judgment calls, anomaly recognition, and overrides under pressure.
- **Human in the lead:** humans move upstream from manual QA of algorithmic output to setting strategic intent, defining ethical guardrails, and proactively steering AI objectives^{[1][4]}.

"Human oversight cannot be a ceremonial checkpoint at the end of an automated process... The goal is not simply human in the loop. It is human in the lead."

— Reggie Townsend, VP Data Ethics, SAS (quoted in source material^[1])

5.3 Cross-functional organizational design

Because AI lifecycle governance touches clinical/business outcomes, legal compliance, risk, and technology, siloed teams fail^[2]. Effective governance requires a **named senior owner** for each material AI deployment and a **cross-functional oversight body with actual authority to pause, modify, or decommission deployments on demand**^{[1][2]}. This is the "intervention node" concept: the authority to pause an autonomous deployment cannot be distributed across five committees on quarterly schedules — it must be concentrated in a cross-functional, real-time body with the explicit technical and political authority to modify or halt systems on demand^[5]. Risk, compliance, and legal specialists should be embedded early in development teams to define minimum risk standards and observe the build — not act as late-stage vetoes^[2].

5.4 Infrastructure-level technical controls

Agentic systems can fail faster than human operators can respond, so hard technical constraints are essential^[2]:

- **Agent identity and least privilege:** treat AI agents as non-human actors with distinct identities and least-privilege credentials, preventing agents from borrowing broad human permissions and ensuring accurate audit trails of exactly what the agent accessed.
- **Kill switches and circuit breakers:** global hard stops revoking tool permissions in seconds, rate governors capping API budgets, and per-agent circuit breakers for repetitive or high-cost actions.
- **Examination-ready documentation:** infrastructure that automatically logs the full decision chain — prompts, plans, tool calls, intermediate results, final actions — so complete audit trails are a byproduct of the system running, not a manual pre-audit assembly exercise^{[1][2]}.
- **Intent-chain telemetry:** regulations ask institutions to explain decisions, and an agent's execution is a *chain* — every prompt, memory state, tool call, intermediate result, and rationale must be stored in tamper-evident audit logs. Logs that capture outputs but not the chain can prove *what* happened, but not *why*^[5]. Production exemplars exist: Cyber Sierra's TracyAI, an AI governance, risk and

compliance (AGRC) agent, generates compliance responses in roughly 15 minutes using multi-step reasoning over past documents with strict access controls and real-time validation^[5].

6. Technical Architecture for Agentic Governance

6.1 The changed risk profile: from generated text to autonomous action

The 2026 Blueprint draws the fundamental distinction^[3]: traditional AI and chatbots perform read-only analysis and content generation, where the primary risks are hallucination, bias, and bad information, controlled through output filters and human proofreading, with low system impact. Agentic AI executes tasks and interacts with software; its primary risks are **irreversible actions, state changes, and corrupted data**; its controls must be sandboxing, branch chains, and system-level architecture; and its system impact is severe — an unbounded agent can delete data or violate regulations. Autonomy amplifies risk: when AI moves from suggesting to doing, governance must shift from content filtering to action bounding^[4].

The contrast can be expressed precisely: legacy model risk frameworks were written for models that produce outputs a human reviews; agents execute^[5]. Under agentic governance, the unit of risk shifts from the model artifact to the **action surface** (tool access, data movement); the analytic focus shifts from model behavior to the **intent chain**; evidence shifts from validation reports to **real-time telemetry**; and validation itself shifts from point-in-time testing to **dynamic guardrails and continuous behavioral testing**^[5].

6.2 The shadow agent threat

Two data points frame the exposure^[4]: **29% of employees** are already using unsanctioned AI agents for work tasks, and **35% of organizations** admit they could not shut down a rogue AI agent if one emerged. Because agents inherit human permissions, an unsanctioned agent with broad tool-calling access is not merely a compliance failure — it is an active operational liability. Countermeasures include living inventories with automated detection, and a Machine Learning Bill of Materials (ML-BoM) logging the exact lineage, training data, and vendor dependencies of every deployed model under first- and second-line-of-defense controls^{[3][4]}.

6.3 The three-layer sandbox

Table 4 The 3-Layer Sandbox for agentic systems^[4]

Perimeter	Function	Key controls
1. Model layer (alignment)	Prevent policy violations before data leaves the enterprise	Prompt guardrails; PII detection and masking; data sanitation

2. Orchestration layer (behavior)	Constrain runaway process behavior	Infinite-loop detection; hard-coded timeouts; mandatory interruptibility; structural kill switch
3. Tool layer (least privilege)	Restrict the agent's blast radius	Role-based access control; restriction to highly specific APIs (e.g., can read inventory but cannot trigger purchase orders)

Complementing the sandbox, the "architect's harness" adds **bounded costs** (the cost of failure must be strictly contained for autonomous automation to work economically), **simulated testing** (because the natural-language input space is nearly infinite, safety must be tested by simulating user behavior, not only at the application layer), and **branch-chain rollbacks** (checkpointing that lets an agent attempt an action, fail safely in a sandbox, and roll back without corrupting production data)^[3].

6.4 Risk-tiered autonomy limits

Autonomy must be designed, not assumed — embedded into product design rather than bolted on post-launch^[4]. The Walmart model illustrates the principle: when an AI system can trigger a purchase or alter inventory, it is engineered with predefined authority boundaries, mapping cleanly to the SR 11-7 "effective challenge" principle so agents cannot independently execute consequential decisions^[4].

Table 5 Risk-tiered autonomy model^[4]

Tier	Autonomy level	Examples	Oversight requirement
1	Full autonomy	Internal productivity, spam filters	AI drafts and executes low-stakes tasks; human applies the output
2	Bounded autonomy	Customer-service chatbots	Resolves standard queries independently; hard-coded escalation above thresholds (e.g., refunds over \$50)
3	Zero autonomy	Credit, HR, life/safety	Explicit human approval structurally required for every action; maps to EU AI Act Annex III and SR 11-7 guidelines

6.5 Calibrating governance to the level of agency

Governance must be calibrated to the level of agency a system actually exercises^[5]. MSD (pharmaceuticals) implemented risk-calibrated pathways for its 75,000 employees based on five levels of agency, scored across dimensions including action scope, permission scope (read-only vs. controlled write vs. full write), initiative (reactive vs. proactive), goal complexity, human-oversight mode, risk impact, and data/context scope. Under this model, low-risk agents (levels 1–2) operate with continuous or periodic oversight for task-specific functions, while high-risk, self-directed agents (levels 4–5) with broad data access require predefined alert and exception-only escalation pathways^[5]. Governance is thus a sliding scale, not a uniform standard: organizations strategically distribute authority based on domain stakes — fintech automates fraud-detection

processes while strictly preserving human decisions; clinical support keeps process autonomy deliberately lower to protect high-stakes human welfare; and the dimensional governance plot (process autonomy vs. decision authority, with circle size representing accountability allocation) provides a practical tool for portfolio-level review^[5].

6.6 The agentic governance control plane

Five components make up the control plane^[4]: (1) a **registry** — a centralized single source of truth for all agents, owners, and risk tiers; (2) **access control** — identity-driven, least-privilege permissions for non-human users operating at machine speed; (3) **visualization** — real-time observability telemetry tracking agent behavior, API dependencies, and data touches; (4) **interoperability** — consistent governance enforcement across open-source and third-party vendor ecosystems; and (5) **security** — continuous posture scanning and anomaly detection to identify threats pre-exfiltration. Frictionless intake matters equally: governance fails if it blocks work, so intake processes must classify systems by minimal, limited, or high risk and apply proportionate controls^[3].

6.7 Board-level governance machinery

At the top of the structure, boards must approach AI with ambition, not fear: robust corporate governance is the structural integrity that lets organizations move confidently and quickly^[3]. The Blueprint assigns the board four workstreams — strategic oversight, active technology & security, workforce transformation, and trustworthy AI — and three structural adaptations^[3]:

- **AI-tailored risk management:** reevaluate existing risk frameworks for the new operational, reputational, and competitive threats AI introduces.
- **Evolving board structures:** revise committee mandates, review frequency, and agenda allocation to match the speed of AI development.
- **Transparent reporting:** demand accessible explanations of major AI capital-allocation decisions, categorized strictly as risks, opportunities, or both.

Strategically, boards must balance short-term productivity pressure against foundational investments (data, infrastructure, talent, even when returns are hard to quantify) and long-term value creation (rethinking value-chain composition and building organizational flexibility to scale up, scale down, and change course as technology evolves)^[3]. On the workforce side, organizations must define which decisions demand meaningful human judgment as non-negotiable, shift HR strategy toward building sustainable AI-augmented environments, and honestly assess the long-term impact on entry-level talent pipelines^[3].

7. Legal Risk Assessment

7.1 Methodology

This section applies a standard severity × likelihood legal risk framework. **Severity** is rated 1 (Negligible) to 5 (Critical) based on potential financial exposure, operational disruption, reputational damage, and regulatory consequences; **likelihood** is rated 1 (Remote) to 5 (Almost Certain) based on precedent, triggering events, and current conditions documented in the source material. The risk score (severity × likelihood) maps to four action bands: **GREEN (1–4, low)**, **YELLOW (5–9, medium)**, **ORANGE (10–15, high)**, and **RED (16–25, critical)**. This assessment supports legal workflows but is not legal advice; it should be reviewed by qualified legal professionals and calibrated to the organization's own risk appetite.

7.1.1 The liability principle: autonomy transfers liability upward, never away

One principle from the source material should anchor every legal analysis of agentic AI: regulators consistently treat the organization as responsible for an agent's actions, **exactly as if an employee had taken them**. If an agent calls a tool, moves data, or initiates a transaction, the institution holds the ultimate accountability for that execution — autonomy transfers liability *upward* to the board and C-suite, never away from the institution^[5]. Vendor contracts, model cards, and disclaimers do not discharge this accountability; they merely create contractual recourse after the fact. This principle is what elevates the technical controls of Sections 5 and 6 — least-privilege agent identity, bounded autonomy, intent-chain telemetry, and kill switches — from engineering best practice to legal necessity, and it underpins the severity ratings below.

7.2 Risk register and scoring

Table 6 AI governance risk register (assessment date: July 31, 2026)

ID	Risk	Category	Severity	Likelihood	Score	Level
R1	Agentic model drift defeats SR 11-7 periodic validation	Regulatory / Model Risk	4 — High	4 — Likely	16	RED
R2	Shadow AI and rogue-agent exposure	Regulatory / Security	4 — High	4 — Likely	16	RED
R3	Algorithmic claim denials causing patient harm	Litigation / Regulatory	5 — Critical	5 — Almost Certain	25	RED
R4	Multi-jurisdiction regulatory fragmentation	Regulatory	3 — Moderate	5 — Almost Certain	15	ORANGE
R5	Vendor / AI supply-chain failure (multi-vendor liability)	Contract / Operational	3 — Moderate	3 — Possible	9	YELLOW
R6	Unbounded autonomous actions (irreversible state changes)	Operational / Regulatory	5 — Critical	3 — Possible	15	ORANGE
R7	Oversight fatigue and automation complacency	Operational / Employment	3 — Moderate	4 — Likely	12	ORANGE
R8	Documentation and examination-readiness gap	Regulatory	4 — High	4 — Likely	16	RED
R9	Data privacy, consent, and PII exposure through agents	Data Privacy	4 — High	3 — Possible	12	ORANGE
R10	Workforce deskilling and hollowed talent pipeline	Employment / Strategic	2 — Low	4 — Likely	8	YELLOW

7.3 Detailed assessments of the top risks

R3 — Algorithmic claim denials causing patient harm (Score 25, RED)

Severity 5 (Critical). Documented evidence shows AI utilization-review systems in Medicare Advantage producing an 82% appeal-overturn rate — i.e., the large majority of contested denials were wrong. Denied care translates directly into patient harm, class-action exposure, regulatory enforcement, and severe reputational damage; officer-level accountability is increasingly plausible under emerging statutes^[1].

Likelihood 5 (Almost Certain). Three in four health plans already use AI for prior authorization, and fewer than 1% of patients appeal — the triggering conditions are present and operating today, not prospective^[1].

Contributing factors: absence of corrective feedback loops; volume-driven automation incentives; FDA capacity constraints limiting external backstop. **Mitigations:** qualified-human sign-off on every adverse determination (anticipating burden-shifting laws such as Louisiana SB 246); compliance with SB 1120-type prohibitions on AI-only medical-necessity determinations; fairness and overturn-rate monitoring as continuous controls^{[1][2]}.

Recommended approach: immediate escalation to General Counsel and the board; engage outside counsel with health-law and class-action expertise; restructure utilization-review workflows now rather than in reaction to legislation. **Residual risk after mitigation:** YELLOW (6–8). **Monitoring:** weekly overturn-rate and appeal-volume dashboards; re-assessment triggered by any new state burden-shifting statute.

R1 — Agentic model drift defeats SR 11-7 periodic validation (Score 16, RED)

Severity 4 (High). A framework failure in model risk management exposes the institution to supervisory criticism, matters-requiring-attention, and enforcement — particularly where agents autonomously touch trade accounting, KYC, and compliance-exception workflows^[1].

Likelihood 4 (Likely). Material behavioral change between formal review events is an inherent property of autonomous recalibration, not a tail event; the 15-point relevance-to-mitigation gap in the 2025 Stanford AI Index confirms most institutions have not closed it^[1]. Gartner's prediction that governance gaps will cause 50% of AI agent deployment failures by 2030 reinforces that this is a sector-wide condition, not an idiosyncratic one^[5].

Mitigation options: tiered validation calibrated to autonomy level; continuous champion-challenger testing; drift detection on input/output distributions embedded in the production stack; runtime policy enforcement at action level^{[1][2]}. **Residual risk:** YELLOW (8). **Next steps:** map all material models to autonomy tiers (CRO, 30 days); stand up continuous drift monitoring for tier-1 agents (CTO, 90 days).

R2 — Shadow AI and rogue-agent exposure (Score 16, RED)

Severity 4 (High). Unsanctioned agents inheriting broad human permissions create data-loss, compliance, and security exposure with no audit trail; the Blueprint treats shadow AI as a first-order legal and security vulnerability^[3].

Likelihood 4 (Likely). 29% of employees already use unsanctioned agents and 35% of organizations could not shut down a rogue agent — prevalence and incapacity are both documented^[4].

Mitigations: living AI inventory with automated endpoint and vendor-integration scanning; centralized agent registry; distinct non-human identities with least-privilege credentials; kill-switch capability as a precondition for production deployment^{[1][4]}. **Residual risk:** YELLOW (6). **Monitoring:** continuous discovery scans; monthly registry reconciliation.

R8 — Documentation and examination-readiness gap (Score 16, RED)

Severity 4 (High). Regulators increasingly evaluate AI maturity through examination-ready artifacts; institutions that manually assemble documentation before audits will fail realtime supervisory expectations and may face findings even where underlying systems are sound^[1].

Likelihood 4 (Likely). With only ~14% of enterprises fully prepared to deploy AI capabilities and governance "advancing on paper" while production infrastructure lags, examination shortfalls are foreseeable across the sector^[3].

Mitigations: automated decision-chain logging (prompts, plans, tool calls, intermediate results, final actions) in tamper-evident logs as a byproduct of system operation; intent-chain telemetry that records not just outputs but the reasoning chain, so the institution can prove *why*, not merely *what*^[5]; ML-BoM lineage records for every deployed model^{[2][3]}. **Residual risk:** GREEN–YELLOW (4–6). **Monitoring:** quarterly mock examinations; log-completeness metrics.

7.4 Escalation and outside-counsel criteria

Consistent with standard escalation practice, the following conditions in this domain warrant defined responses:

- **Immediate board/GC escalation (RED items):** R1, R2, R3, R8. Establish a dedicated response team per item, define contingency plans, and review weekly until scores fall below 16.
- **Mandatory outside-counsel engagement:** any active litigation over algorithmic denials (R3), any regulatory inquiry or examination finding (R1, R8), and any government investigation touching AI deployments.
- **Strongly recommended engagement:** novel questions of first impression (liability allocation for multi-vendor agent failures, R5); jurisdictional complexity from conflicting state AI statutes (R4);

new-regulation compliance program build-out (EU AI Act, Vietnam, ASEAN milestones).

- **Consider engagement:** insurance-coverage review for AI-related claims; employment counsel for workforce-transformation impacts (R10); data-incident counsel for agent-mediated PII exposure (R9).

8. Implementation Roadmap

The source material provides a convergent, phased path from boardroom strategy to operational machinery.

8.1 The 90-day playbook

Table 7 90-day implementation sequence^[4]

Phase	Focus	Key actions
Days 1–30	Inventory & evaluations	Spin up an evaluation harness (e.g., UK AISI's <i>inspect_ai</i>); establish the centralized registry; baseline active and shadow agents
Days 31–60	Sandboxing & workflows	Install a reference agent in a sandbox; stand up the cross-functional AI Governance Council; define risk tiers, autonomy limits, and approval workflows
Days 61–90	Continuous observability	Integrate vulnerability scanning (e.g., NVIDIA's <i>garak</i>) and real-time telemetry tied to business KPIs; enforce structural kill-switch capability for all production agents

8.2 Strategic build-out (quarters 2–4)

- **Governance machinery:** adapt board committee mandates and reporting; map KPMG pillar ownership to NIST Govern–Map–Measure–Manage functions^[3].
- **Human oversight redesign:** move reviewers from end-of-pipe checking to simulator-trained, risk-routed "human-in-the-lead" roles; define non-negotiable human-judgment decision classes^{[2][3]}.
- **Sector tracks:** financial institutions prioritize living inventory, tiered validation, and circuit breakers; healthcare institutions prioritize multi-state disclosure infrastructure, physician-level review of adverse determinations, and procurement-stage vendor due diligence^[1].
- **Legal workstream:** register all Section 7 risks in the legal risk register with owners and review dates; execute the escalation plan for RED items; pre-brief outside counsel for the mandatory-engagement triggers.

9. Conclusion

The transition from theoretical oversight to operational readiness requires structural integrity at every level: boardroom strategy, NIST-style risk machinery, and architectural harnesses bounding agentic AI^[3]. The evidence across all four sources points the same way: governance is advancing on paper while production-

ready infrastructure remains critically underdeveloped^[3]; regulators in every major jurisdiction are converging on lifecycle accountability^[1]; and the institutions closing the gap treat governance as continuous, runtime, cross-functional organizational design rather than periodic legal compliance^[2].

The legal risk posture is equally clear: the gravest exposures — algorithmic denial of care, uncontrolled agentic drift, shadow agents, and documentation gaps — are already materializing, not hypothetical, and liability for them transfers upward to the institution, never away^[5]. They are also governable. Governance designed with humans placed deliberately rather than defensively converts "friction" from a failure into a feature — resistance in the deployment flow is the signal that AI has reached a level of operational importance where trust and accountability must evolve alongside capability^[5]. Organizations that embed agentic governance early do not merely reduce operational risk; they move faster, because governance built correctly is not a brake on innovation but the fuel that lets it scale^[4]. The framework for accountable autonomy is visible; the only remaining question is whether an institution's governance can act as fast as its agents do^[5].

"When everyone is one prompt away from the same answers, human judgment becomes the differentiator... The strongest AI strategies use automation to scale human expertise, not eliminate it."

— Bryan Harris, CTO, SAS (quoted in source material^[1])

This report synthesizes the user-supplied source documents and curated source library listed in the References and Appendix, and applies a generic legal risk assessment framework. It does not constitute legal advice. Risk ratings reflect the evidence in the source material and should be reviewed by qualified legal professionals and calibrated to the organization's specific risk appetite, jurisdictions, and industry context.

References

1. *AI Governance in Regulated Industries: Navigating the Lifecycle Accountability Gap* (research brief, Gemini Notebook; user-supplied document, 2026).
2. *AI Governance Strategies in Highly Regulated Industries* (analysis of governance practices in financial services and healthcare; user-supplied document).
3. *The 2026 Blueprint for Enterprise AI Governance: A Strategic Guide for Boards and Risk Officers* (synthesizing KPMG, NIST, and Columbia Data Science Institute insights; user-supplied and authored analysis).
4. *Governing the Agentic Enterprise: Practically Setting Autonomy Boundaries and Evolving Human Oversight from In-the-Loop to In-the-Lead* (user-supplied and authored analysis).
5. *The Accountability Gap: How Highly Regulated Industries are Governing Agentic AI* (user-supplied and authored analysis).

6. *Source List: Governing Agentic AI* (curated bibliography of regulations, frameworks, academic research, and industry guidance; user-supplied PDF — reproduced in the Appendix).

Appendix: Curated Source Library

The following source list, supplied with the briefing materials, catalogues the regulations, standards, research, and industry guidance underpinning the governance practices described in this report. It is reproduced here, categorized by topic and source type, as a working reference for legal and compliance teams.

A. Regulations, Frameworks & Standards

- Article 14: Human Oversight — EU Artificial Intelligence Act
- Concept Note: AI Risk Management Framework — Trustworthy AI in Critical Infrastructure Profile — NIST
- EU AI Act Compliance Guide: Risk & Obligations — HCLTech
- ISO/IEC 42001 Certification: The Global Standard for AI Management Systems (AIMS) — KPMG
- Japan's Second Artificial Intelligence Basic Plan
- Model AI Governance Framework for Agentic AI — IMDA (Singapore)
- NIST AI Risk Management Framework (AI RMF 1.0)
- Overview of ISO/IEC 42001 (AI Management System) and AI System Conformity Assessment — UNIDO
- Reporting Framework — Hiroshima AI Process
- Safeguards for Agentic Finance at Runtime — Monetary Authority of Singapore
- US FDA Artificial Intelligence and Machine Learning Discussion Paper

B. Academic & Research Papers

- A Decoupled Human-in-the-Loop System for Controlled Autonomy in Agentic Workflows — arXiv
- AI Safety and Automation Bias — CSET
- AI-Instigated Human Oversight: Rethinking Human-in-the-Loop Safety in Clinical AI
- Constitutive vs. Corrective: A Causal Taxonomy of Human Runtime Involvement in AI Systems — arXiv
- FRAME: Comprehensive Risk Assessment Framework for Adversarial Machine Learning Threats — arXiv
- Governance Architecture for Autonomous Agent Systems: Threats, Framework, and Engineering Practice — arXiv
- Human-AI Governance (HAIG): A Trust-Utility Approach — arXiv
- Insurance of Agentic AI — arXiv

- Metis AI: The Overlooked Middle Zone Between AI-Native and World-Movers — arXiv

C. Industry Reports & Corporate Guidelines

- Agentic AI Governance Standards and Regulations: What Actually Applies in 2026
- Agentic AI Governance, Risk, and Controls for Business Leaders — Bain & Company
- Agentic AI Governance: NIST Standards for Autonomous Systems — Cloud Security Alliance
- AI Ethics Board Structure Guide — VerifyWise AI Governance Library
- AI Governance in Regulated Industries — Horizon Scan 001
- AI Leadership: Turning Investment into Value — Accenture
- AI Red Teaming: The Complete Guide for Security Teams (2026) — Repello AI
- AI Leaders Podcast Ep. 80: Human in the Lead — Redesigning Agency with Trust and AI (transcript) — Accenture
- Best Practices for AI in Regulated Industries — Jama Software
- Bridging the Principle–Practice Gap in Responsible AI: A Cross-Domain Review
- How to Build Human-in-the-Loop Oversight for Production AI Agents — Galileo AI
- Human in the Lead, Not Human in the Loop: Rethinking AI Governance — Kategos AI
- Human-in-the-Loop AI Workflows That Actually Scale — Mak it Solutions
- Human-in-the-Loop vs Human-on-the-Loop in Agentic AI — Tek Leaders
- Human-in-the-Loop: A 2026 Guide to AI Oversight That Actually Works — Strata Identity
- KPMG Make AI Scale: From Experimentation to Transformation
- Model Risk Management in the Age of AI — Moody's
- NIST AI RMF Human Oversight Controls: A Practical Guide — Living Security
- SAS: For Agentic AI ROI, Invest in Human Judgment
- Speech by Vice Chair for Supervision Bowman on Artificial Intelligence in the Financial System — Federal Reserve
- The Agentic Oversight Framework: Procedures, Accountability — Sardine AI
- The Challengers: Ethics, Technology and the Evolving Role of the CFO — CPA Ontario
- The Difference Between AI Risk Management and AI Governance — Airia
- The Multi-Year AI Advantage — Capgemini
- Third-Party AI Risk Management: Managing Vendor and Model Risk — Unorma
- What Is a Chief AI Officer? — IBM

D. General Reference

- Artificial Intelligence Act — Wikipedia