

EXECUTIVE BRIEFING

Strategic Briefing: Governance & Control Architecture for Agentic AI



Navigating the Supervisory Expectations Gap Post-SR 26-2

PREPARED FOR THE AUDIT COMMITTEE & EXECUTIVE RISK BOARD

Synopra Strategic Insights

The Core Exposure Gap

THE SUPERVISORY REALITY IN MID-2026



The Policy Mirage: Traditional Model Risk Management (SR 11-7) structures are built for static statistical models, not autonomous loops.



The SR 26-2 Exclusions: Recent updates leave generative and agentic workflows technically outside standard MRM boundaries, yet examiners still evaluate operational controls.



The Risk: Autonomous agents are already interacting with production databases, code bases, and customer-facing channels without an auditable verification trail.

The 5-Stage Agent Control Framework

Moving From Paper Policies to Running Controls

1

1. Inventory

Establishing a programmatic discovery mechanism for both sanctioned and shadow enterprise agents.

2

2. Classify

Tiering agent use cases by materiality and deterministic risk profiles.

3

3. Control

Embedding human-in-the-loop gates, scope boundaries, and API kill switches directly inside the automation stack.

4

4. Evidence

Maintaining a real-time log of prompts, tool responses, and systemic decisions that can withstand second-line testing.

5

5. Oversee

Establishing continuous drift and semantic monitoring rhythms for the Board.

Key Board Inquiries to Anticipate

THREE QUESTIONS EXAMINERS AND BOARDS WILL ASK THIS QUARTER

01

Where are our agents executing, what **data boundaries** are enforced, and how do we actively prevent shadow deployments?

02

If an autonomous agent hallucinates a transaction or data-deletion loop, what **real-time programmatic fallback** or kill switch prevents enterprise-wide contagion?

03

Can we produce an **immutable, itemized audit log** of every step, prompt, and tool call an agent executed during a disputed customer transaction?

The Tactical Roadmap

Transitioning Frameworks into Running Infrastructure



Weeks 1–3

Complete a comprehensive cross-departmental AI Agent Exposure & Risk Classification inventory.



Days 30–90

Select one high-priority agent workflow to re-architect with embedded verification gates and evidence pipelines.



Ongoing

Implement fractional AI risk leadership rhythms to validate controls before systemic examination.

STRATEGIC OBJECTIVE

Transitioning Policy Frameworks into resilient Audit-Ready Infrastructure by Q4 2026 to meet evolving regulatory expectations.

Hyperscaler Compliance Integration

Synopra advisory and implementation services bridge native cloud security with SR 26-2 agentic AI control requirements.



AWS Compliance Foundation

Leveraging an extensive portfolio of security standards and certifications — including NIST, ISO, SOC, HIPAA, GDPR, and PCI DSS — to secure the underlying compute and data layers.



Azure Compliance Foundation

Building on Azure's broad compliance portfolio across global (ISO/SOC), U.S. Government (FedRAMP High, DoD, ITAR), industry (HITRUST, FDA), and regional frameworks.



Synopra SR 26-2 Overlay

Our implementation services map these natively secure cloud primitives directly to the operational control, logging, and human-in-the-loop requirements specific to autonomous agent governance.

Turn Frameworks into Defensible Controls

Synopra helps regulated institutions make agentic AI examiner-ready — with running controls, audit trails, and evidence, not policy binders.

Start with a 30-minute governance briefing

synopra.com · martin@agilentic.com

This briefing is provided for general information only and does not constitute legal, regulatory, or compliance advice. Regulatory references reflect our understanding at the time of publication and may change. Organizations remain responsible for their own compliance obligations.