

AI Governance for Financial Services

Model Risk Management and Regulatory Compliance

A guide for SR 26-2 model risk management, SEC and FINRA compliance, and enterprise AI risk management.

Governance is enforced at the authorization boundary: the point where an output becomes an effect-bearing action, and where evidence, not explanations, must exist.

PUBLISHER

FERZ, Inc.

VERSION

1.1 · June 2026

LICENSE

CC BY 4.0

FOR

Chief Risk Officers · Chief Model Risk Officers · Heads of AI Governance · Compliance Officers · Internal Audit ·
Model Validation Teams

Front Matter

License and Terms of Use

Copyright and License

Copyright © 2026 FERZ, Inc. All rights reserved.

This work is licensed under the Creative Commons Attribution 4.0 International License (CC BY 4.0). You are free to share (copy and redistribute the material in any medium or format) and adapt (remix, transform, and build upon the material for any purpose, including commercial purposes) under the following terms:

Attribution: You must give appropriate credit, provide a link to the license, and indicate if changes were made. Attribution must not suggest endorsement by the licensor.

The full license text is available at creativecommons.org/licenses/by/4.0.

Attribution Guidance

When citing or reusing this work, attribution must include:

FERZ, Inc., "AI Governance Executive Guide: Financial Services," 2026. Available at ferz.ai.

If the material is modified, adapted, or excerpted, the attribution must clearly indicate that changes were made.

Commercial Scope Notice

This publication is provided as a public reference framework for AI governance and risk management. While reuse and adaptation are permitted under the CC BY 4.0 license, commercial services, tooling, certification programs, managed governance services, and implementation accelerators derived from this framework are offered separately by FERZ, Inc. or its authorized partners. This notice does not restrict rights granted under the license; it clarifies the separation between open reference material and commercial offerings.

Framework Scope and Citation Currency

This guide is a foundational governance framework, not an exhaustive policy tracker or current-law reference. It cites statutes, regulations, memoranda, standards, and guidance only where they materially affect the general governance model, compliance posture, evidence requirements, oversight structure, or audit-readiness framework. Use-case-specific, agency-specific, jurisdiction-specific, or rapidly changing requirements may apply and should be verified separately before reliance or implementation.

No Warranty or Endorsement

This document is provided "as is," without warranties of any kind, express or implied. Use of this framework does not constitute regulatory, legal, or compliance advice, nor does it imply endorsement by FERZ, Inc. of any implementation, product, or service that references it.

Limitation of Liability

To the maximum extent permitted by law, FERZ, Inc. is not liable for any direct, indirect, incidental, consequential, special, exemplary, or other damages arising from use of, reliance on, adaptation of, or implementation based on this publication. Users are responsible for obtaining appropriate legal, regulatory, procurement, technical, and compliance review before applying any framework, checklist, clause language, or procedure described in this guide.

Canonical Reference

The authoritative version of this document is maintained by FERZ, Inc. Subsequent adaptations or derivatives should reference the original source to ensure conceptual fidelity and version traceability.

Overview

Executive At-a-Glance

Audience

Chief Risk Officers, Chief Model Risk Officers, Heads of AI Governance, Compliance Officers, Internal Audit, and Model Validation Teams at banks, broker-dealers, investment advisers, and insurance companies.

Core Thesis

Governance is enforced at the authorization boundary: the point where an output becomes an effect-bearing action, and where evidence, not explanations, must exist.

Two Governance Tracks

TRACK	PRIMARY RISK	OWNERSHIP
Decisioning AI	Model risk (SR 26-2, ECOA, Reg BI)	MRM / Model Risk Committee
Communications AI	Records and supervision (17a-4, FINRA 3110)	Compliance / Supervision

Four-Tier Classification

TIER	RISK LEVEL	GOVERNANCE INTENSITY
Tier 1	Critical	Full MRM: validation, bias testing, quarterly review, MRC approval
Tier 2	Elevated	Proportionate MRM: validation, annual review
Tier 3	Standard	Operational controls, data security
Tier 4	Minimal	Acceptable use policy only

90-Day Implementation

Weeks 1 to 4 (Foundation): Inventory sweep, gap assessment, tiering criteria, governance structure.

Weeks 5 to 8 (Critical Systems): Tier 1 focus, SR 26-2 gap analysis, bias testing, validation approach.

Weeks 9 to 12 (Operationalization): Policy deployment, training, Examiner Pack compilation, mock examination.

Maturity Endpoint

Level 5: Independently Verifiable. The organization is designed to produce tamper-evident authorization artifacts for consequential AI actions, with evidence sufficient for independent replay.

Summary

Executive Summary

Financial services institutions face a convergence challenge: quantitative AI and machine learning models that meet the model definition fall under Model Risk Management (MRM), while generative and agentic AI sit outside it, requiring adjacent controls and a clear authorization boundary where outputs become effect-bearing actions.

The 2011 SR 11-7 framework was written for credit, valuation, and risk models. In April 2026 the agencies replaced it with SR 26-2, a risk-based framework scaled to each institution's model risk profile. The revised guidance modernizes model risk management, but it also draws a sharp line: generative AI and agentic AI are outside the scope as novel and rapidly evolving, with institutions directed to apply existing risk principles to the tools the guidance does not cover. That line is the opportunity. Quantitative models belong in the model risk management track; the generative, agentic, and effect-bearing layer needs adjacent controls and, where actions execute, a pre-execution authorization boundary.

This creates both risk and opportunity. Organizations that proactively extend their MRM frameworks to address AI governance gain regulatory credibility, reduce examination risk, and establish operational advantages. Those that wait face increasing examination scrutiny, potential enforcement actions, and competitive disadvantage.

Two Governance Tracks

AI governance in financial services operates across two distinct tracks, each with different regulatory drivers and organizational ownership:

TRACK	PRIMARY RISK	KEY REGULATIONS	TYPICAL OWNERSHIP
Decisioning AI	Model risk	SR 26-2, ECOA, Reg BI	MRM / Model Risk Committee
Communications AI	Records / supervision	SEC 17a-4, FINRA 3110, 3120	Compliance / Supervision

Decisioning AI: Underwriting, AML prioritization, pricing, trading signals, suitability scoring, portfolio optimization. Governed primarily through MRM frameworks with validation, bias testing, and model documentation.

Communications AI: Client-facing drafts, advisor copilots, marketing content, complaint handling, call summarization, research synthesis. Governed primarily through supervision, recordkeeping, and disclosure frameworks.

While communications AI may not directly execute decisions, its regulatory risk profile can be equal or higher due to books-and-records, supervision, disclosure, and misrepresentation exposure. Do not assume that communications use cases are inherently lower risk than decisioning use cases.

This distinction matters because examiners ask different questions, and internal ownership often differs. A single AI system may span both tracks, an advisor copilot that drafts client communications (supervision) based on portfolio analysis (MRM), requiring coordinated governance.

The Core Challenge

Model Risk Management has three pillars: Development, Validation, and Governance. Each requires extension for AI:

PILLAR	TRADITIONAL MODELS	AI / ML EXTENSION REQUIRED
Development	Mathematical specification, assumptions documented	Training data provenance, architecture selection, hyperparameter rationale, explainability approach
Validation	Back-testing, sensitivity analysis, stress testing	Bias and fairness testing, drift monitoring, adversarial testing, explainability assessment
Governance	Model inventory, tiering, periodic review	Expanded inventory (embedded AI, third-party), new risk categories, continuous monitoring

What This Guide Provides

Regulatory mapping: How SR 26-2, SEC and FINRA guidance, fair lending requirements, and the EU AI Act apply to AI systems.

Risk framework: Four-tier classification system with explicit Tier 1 triggers to prevent governance arbitrage.

Governance structure: Three Lines of Defense model extended for AI oversight.

Lifecycle management: Development, validation, deployment, monitoring, and retirement processes.

Examination readiness: Documentation standards, Examiner Pack template, and response playbooks.

GenAI controls: Dedicated control set for generative AI including prompt injection, supervision, and disclosure.

Board metrics: KRIs and KPIs for executive AI governance reporting.

Implementation roadmap: 90-day plan with maturity progression to independently verifiable governance.

90-Day Implementation Overview

PHASE	KEY DELIVERABLES
Weeks 1 to 4: Foundation	AI system inventory sweep, MRM framework gap assessment, tiering criteria established, core policies drafted, governance structure defined
Weeks 5 to 8: Critical Systems	Tier 1 systems identified and prioritized, gap analysis against SR 26-2, bias testing for credit models, validation approach documented, monitoring capabilities established
Weeks 9 to 12: Operationalization	Policies approved and deployed, training completed, shadow AI detection active, Tier 1 documentation packages compiled, mock examination conducted

PART I

The Regulatory Landscape

CHAPTER 1

Existing Frameworks That Apply Today

Financial institutions do not operate in a regulatory vacuum when deploying AI. Multiple existing frameworks apply, even in the absence of AI-specific regulation. Understanding these frameworks, and how regulators are extending them to AI, is essential for defensible governance.

SR 26-2 / OCC 2026-13: The Current U.S. Banking Anchor

The 2011 SR 11-7 framework is now legacy guidance. On April 17, 2026, the Federal Reserve, OCC, and FDIC issued revised model risk management guidance, SR 26-2, OCC Bulletin 2026-13, and FDIC FIL-15-2026, which supersedes and replaces SR 11-7. The revised guidance emphasizes risk-based model governance scaled to each institution's model risk profile, model development and use, validation and ongoing monitoring, governance and controls, and third-party model considerations. For historical continuity, many institutions will still refer to "SR 11-7-style" model risk management, but the current reference point is the 2026 revised guidance, which is most relevant to banking organizations with more than \$30 billion in total assets.

Model definition, narrowed. The revised guidance defines a model as a complex quantitative method with a theoretical underpinning that processes input data into quantitative output. It excludes simple arithmetic calculations such as spreadsheets, as well as deterministic rule-based processes and software. Quantitative AI and machine learning systems that meet this definition belong in the model risk management track; tools that do not are governed through adjacent controls.

Generative and agentic AI are out of scope. The revised guidance states that generative AI and agentic AI models are novel and rapidly evolving and are therefore not within its scope, while directing institutions to apply existing risk management principles to tools and systems it does not cover. Model risk management governs the quantitative models that meet the model definition. Generative AI, agentic tools, communications AI, and effect-bearing workflow automation may require adjacent controls: supervision, recordkeeping, vendor oversight, information-security controls, and, where actions may execute, pre-execution authorization through a runtime authorization boundary. Model risk management governs models; the authorization boundary governs the point where outputs become effect-bearing actions.

The Three Pillars of Model Risk Management

1. Model Development and Implementation: Requires disciplined and knowledgeable development and implementation processes. For AI, this extends to training data selection, architecture choices, hyperparameter tuning, and feature engineering documentation.

2. Model Validation: Requires independent review to verify models perform as expected. For AI, validation must address explainability, bias detection, robustness to adversarial inputs, and performance

under distribution shift.

3. Governance, Policies, and Controls: Requires clear accountability, comprehensive inventory, and risk-based oversight. For AI, governance must expand to cover embedded AI features in vendor systems, third-party AI services, and rapid model iteration cycles.

“Effective Challenge” for AI

The revised guidance retains “effective challenge” as a core principle: critical analysis by objective, informed parties who can identify model limitations. For AI systems, effective challenge requires:

- Technical expertise in AI and ML beyond traditional quantitative modeling.
- Understanding of training data risks (bias, provenance, currency).
- Ability to assess explainability and interpretability approaches.
- Familiarity with AI-specific failure modes (hallucination, drift, adversarial manipulation).

SEC and FINRA: Investment Management and Supervision

For broker-dealers and investment advisers, SEC and FINRA oversight adds requirements that span both decisioning and communications AI.

Fiduciary Duty and Reg BI

Investment advisers have a fiduciary duty to act in clients' best interests. Broker-dealers must comply with Regulation Best Interest (Reg BI). When AI systems make or influence investment recommendations, these duties attach to the AI outputs.

Supervision and Recordkeeping (Communications AI)

AI-generated client communications trigger books-and-records obligations under SEC Rule 17a-4 and FINRA Rules 3110 and 3120. This includes:

- Retention of AI-generated drafts and final communications.
- Supervision of AI-assisted advisor communications.
- Review procedures for AI-generated marketing content.
- Disclosure of AI use in client-facing materials where material.

AI Washing Enforcement

The SEC has brought settled enforcement actions against investment advisers for false or misleading statements about their use of AI. In March 2024, the SEC settled charges against two investment advisers for misrepresenting their use of AI.

Fair Lending and Consumer Protection

AI systems used in credit decisions face heightened scrutiny under fair lending laws. The Equal Credit Opportunity Act (ECOA) and Fair Housing Act prohibit discrimination, regardless of whether discrimination is intentional or results from seemingly neutral AI systems.

CRITICAL REQUIREMENT

Under ECOA's adverse action notice requirements (Regulation B), creditors must provide specific reasons for adverse actions. When AI drives credit decisions, firms must be able to explain the factors that led to adverse outcomes, a direct explainability requirement.

CHAPTER 2

What's Coming

The regulatory landscape for AI in financial services is evolving rapidly. Institutions should anticipate and prepare for emerging requirements rather than waiting for final rules.

EU AI Act: Extraterritorial Impact

The EU AI Act applies in phases and has extraterritorial reach affecting U.S. financial institutions that serve EU customers or deploy AI systems whose outputs are used in the EU.

High-risk classification: Credit scoring and creditworthiness assessment AI is explicitly classified as high-risk, requiring conformity assessments, risk management systems, data governance, technical documentation, human oversight, accuracy and robustness requirements, and registration in an EU database.

Anticipated U.S. Regulatory Evolution

While the specific regulatory trajectory depends on administration priorities, several developments are likely:

- **NIST AI RMF adoption:** The NIST AI Risk Management Framework is increasingly referenced in regulatory guidance and may become a de facto standard.
- **Generative and agentic AI governance:** The agencies excluded generative and agentic AI from the 2026 model risk guidance because those systems are novel and rapidly evolving. Institutions should expect supervisory attention to continue through existing risk-management, vendor-oversight, information-security, consumer-protection, and examination channels while more specific expectations develop.
- **State-level regulation:** States like Colorado have enacted AI-specific laws affecting insurance and credit decisions.
- **Examination intensification:** Regardless of new rulemaking, examination scrutiny of AI systems will continue to increase.

Importantly, supervisory expectations often evolve through examination practice and enforcement actions before formal rulemaking occurs. Institutions should not wait for explicit AI regulation to implement governance frameworks.

PART II

Risk Framework for Financial Services AI

CHAPTER 3

AI System Classification

Expanding the Model Inventory

The first challenge in AI governance is defining what counts as an “AI system” subject to governance. Traditional model inventories focused on internally developed quantitative models. AI governance requires expanding scope significantly.

For governance purposes, any system that materially influences decisions, recommendations, or regulated communications is brought into governance scope, whether as a model under the revised guidance or through adjacent controls, regardless of technical implementation. This avoids definitional debates that delay governance while risk accumulates.

What counts as an AI system in financial services:

- **Traditional ML models:** Credit scoring, fraud detection, AML transaction monitoring, market risk models.
- **Generative AI:** Document drafting, customer communication, research summarization, code generation.
- **Embedded AI features:** AI capabilities within vendor systems (CRM AI assistants, ERP optimization, trading platform analytics).
- **Third-party AI services:** APIs, cloud ML services, outsourced model development.
- **RPA with AI components:** Robotic process automation incorporating machine learning or NLP.

The Four-Tier Framework

Financial services AI governance requires a four-tier system that calibrates oversight intensity to risk. This framework aligns with the revised guidance's emphasis on calibrating model risk management effort to model materiality, complexity, and each institution's risk profile.

TIER	CRITERIA	EXAMPLES	REGULATORY TOUCHPOINTS
Tier 1 (Critical)	Direct customer impact, credit and investment decisions, regulatory reporting	Credit underwriting, algo trading, loan pricing, AML alert prioritization, robo-advisory	SR 26-2 full scope, fair lending (ECOA), Reg BI, EU AI Act high-risk
Tier 2 (Elevated)	Material business decisions, customer-facing with human review	Fraud scoring (with review), customer service assist, compliance surveillance	SR 26-2 proportionate, operational risk, vendor risk management

TIER	CRITERIA	EXAMPLES	REGULATORY TOUCHPOINTS
Tier 3 (Standard)	Internal operations, no direct customer impact	Document summarization, internal search, meeting transcription	Operational controls, data security, acceptable use
Tier 4 (Minimal)	Personal productivity, no sensitive data	General research, formatting assistance, brainstorming	Acceptable use policy only

Automatic Tier 1 Triggers

Tiering frameworks fail when business units “tier down” to avoid validation overhead. To prevent governance arbitrage, the following conditions automatically classify a system as Tier 1, regardless of other assessments:

- 1. Impacts credit decisions, suitability determinations, best interest assessments, pricing, or execution quality.
- 2. Drives or materially influences regulatory reporting (Call Reports, FR Y-9C, Form PF, and similar).
- 3. Produces outputs that become customer-facing determinations (even if “human reviewed”).
- 4. Involves protected class data or proxies in decisioning logic.
- 5. Feeds into capital, liquidity, or stress testing calculations.

Third-Party AI Risk

Third-party AI presents unique governance challenges. The revised guidance retains third-party model accountability: institutions remain responsible for the outputs and limitations of vendor models, and use of vendor products does not diminish the responsibility of the board and senior management to ensure a model meets the institution's standards.

Vendor Minimum Evidence Package

To make third-party oversight actionable, require the following evidence from vendors, scaled by tier:

EVIDENCE REQUIRED	TIER 1	TIER 2	TIER 3
Model card / capability limitations	Required	Required	Requested
Training data governance statement	Required	Required	N/A
Change management and versioning cadence	Required	Required	Requested
Independent testing attestations (bias, robustness)	Required	Risk-based	N/A
Incident history and SLA for remediation	Required	Required	Requested
Audit and inspection cooperation language	Required	Required	Recommended

This is the difference between “we have vendor risk policy” and “we can survive an examination.”

CHAPTER 4

AI-Specific Risk Categories

AI systems present risk categories that extend beyond traditional model risk. Effective governance requires understanding and addressing each category.

Model Risk (Extended)

Training Data Risk: AI model quality depends fundamentally on training data. Risks include biased data leading to discriminatory outputs, incomplete data leading to poor performance on edge cases, outdated data leading to irrelevant predictions, and data provenance uncertainty affecting auditability.

Architecture Risk: The choice of model architecture affects performance, explainability, and failure modes.

Drift Risk: AI models can degrade over time as the relationship between inputs and outcomes changes. Concept drift and data drift both require detection and remediation.

Explainability Risk: Complex AI models may produce accurate predictions without interpretable reasoning. For regulated decisions requiring explanation, this creates compliance risk.

Operational Risk

Integration failures: AI systems failing to properly integrate with downstream business processes.

Latency and availability: AI systems not meeting performance requirements for real-time decisions.

Rollback capabilities: Inability to quickly revert to prior model versions when issues arise.

Compliance Risk

Fair lending violations: Discriminatory outcomes in credit decisions, even unintentional.

Disclosure failures: Inadequate disclosure of AI use to customers or regulators.

Fiduciary breaches: AI recommendations that do not serve client interests.

Cybersecurity and Data Risk

Training data poisoning: Adversaries corrupting training data to influence model behavior.

Adversarial attacks: Inputs crafted to cause model misclassification or errors.

Data leakage: Sensitive information exposed through AI outputs or model inversion.

Prompt injection: Manipulating AI systems through crafted inputs to bypass controls.

CHAPTER 5

Generative AI Controls

Generative AI presents unique control requirements beyond traditional ML governance. Because GenAI can produce novel outputs (text, code, images) that may be used externally or influence decisions, a dedicated control set is required.

Data Security Controls

- **Input filtering:** Prevent sensitive data (PII, NPI, material non-public information) from entering GenAI prompts.
- **Output monitoring:** Detect and block potential data leakage in AI-generated content.
- **Prompt injection defense:** Technical controls against manipulation attempts that could bypass safety rails or extract training data.
- **Session isolation:** Ensure user sessions do not leak information across conversations.

Supervision and Recordkeeping Controls

For AI-generated client communications, supervision requirements under FINRA 3110 and 3120 and SEC 17a-4 apply:

- **Draft retention:** Capture AI-generated drafts as well as final sent communications.
- **Pre-send review:** Define which communications require human review before sending.
- **Supervision sampling:** Include AI-generated content in routine supervision review sampling.
- **Disclosure requirements:** Determine when AI use requires disclosure to clients.

Use Case Governance

CATEGORY	EXAMPLES	REQUIRED CONTROLS
Approved	Research summarization, internal drafting, code assist	Acceptable use acknowledgment, data classification
Restricted	Client communication drafts, marketing content, regulatory filings	Human review, supervision capture, approval workflow
Prohibited	Automated trading signals, unreviewed client communications, credit decisioning	Technical enforcement, exception process only

Citation and Source Attribution

For research summarization and analysis use cases, require:

- Clear indication when content is AI-generated versus human-authored.
- Source attribution for factual claims where underlying sources exist.
- Confidence indicators where the AI is synthesizing or inferring.
- Prohibition on AI-generated content being presented as proprietary research.

Human Approval Requirements

Define approval gates for external GenAI outputs:

- All client-facing communications require human review and explicit send action.
- Marketing and advertising content requires compliance approval.
- Regulatory correspondence requires legal sign-off.
- Social media posts require designated reviewer approval.

MEANINGFUL REVIEW

Human review must be meaningful, not rubber-stamp. Reviewers must have the authority, training, and time to override or block AI outputs. Checkbox acknowledgments without substantive review do not satisfy human oversight requirements under model risk management or EU AI Act standards.

PART III

Governance Structure

CHAPTER 6

Organizational Model

Three Lines of Defense for AI

Financial services institutions operate under a three lines of defense model. AI governance must integrate into this structure rather than creating parallel governance.

LINE	TRADITIONAL ROLE	AI GOVERNANCE EXTENSION
First Line (Business Units)	Own and manage risk in daily operations	AI system ownership, use case documentation, performance monitoring, issue escalation, data quality for AI inputs
Second Line (MRM / AI Gov)	Set policy, monitor risk, independent validation	AI policy and standards, independent AI validation, expanded inventory management, regulatory liaison, AI-specific risk monitoring
Third Line (Internal Audit)	Independent assurance of framework effectiveness	AI governance framework effectiveness, control testing for AI systems, regulatory compliance verification, AI-specific audit techniques

Key Roles

Chief Model Risk Officer (Expanded Scope): Traditional CMRO role expanded to include AI and ML systems. Requires augmented team with AI expertise.

AI Ethics Officer / Responsible AI Lead: Dedicated focus on AI-specific considerations, including bias, fairness, transparency, and ethical use.

AI Validation Team: Specialized validation capability with skills in ML testing, bias detection, explainability assessment, and adversarial robustness testing.

Business Unit AI Champions: First-line contacts who understand both business context and AI capabilities.

Governance Bodies

Model Risk Committee (Expanded Charter): Existing MRC charter should be amended to explicitly include AI and ML systems. Committee reviews Tier 1 AI deployments that fall within the MRM track and coordinates with AI governance / compliance owners for adjacent-control systems.

AI Ethics Review Board: New body or MRC subcommittee addressing AI-specific ethical considerations.

Executive AI Steering Committee: Senior leadership oversight of AI strategy, risk appetite, and major investments.

GOVERNANCE AUTHORITY

Authority and Capability Requirements

AI governance frameworks fail in two predictable ways: the governance function lacks authority to enforce standards, or it lacks competence to make credible decisions. Both failure modes must be addressed by design.

The Authority Problem

Governance without authority produces policy documents that IT ignores and business units route around. Effective AI governance requires:

- **Pipeline integration:** Governance must have hooks into actual deployment processes, intake forms that trigger before procurement, validation gates that block production deployment, monitoring access with real data. Policy without pipeline integration is theater.
- **Escalation path with teeth:** When governance raises concerns, there must be an escalation path to someone who can override IT delivery pressure and business revenue pressure. This typically means reporting to the CRO, a Board committee, or having direct access to the CEO on material issues.
- **Information rights:** Even where governance cannot block deployments, it must have visibility into all AI systems. Shadow AI is a governance design failure, not just a user behavior problem. Information rights include access to IT asset inventories, procurement notifications, vendor contract reviews, and production deployment logs.
- **Budget independence:** Governance cannot depend on the goodwill of functions it oversees. Validation resources, monitoring tools, and external expertise must be funded independently of IT or business unit budgets.

The Competence Problem

Authority without competence produces either rubber-stamp approvals (governance theater) or reflexive blocking (innovation killing). Neither serves the organization. Effective AI governance requires dual competency:

- **Regulatory and risk expertise:** Deep understanding of model risk management guidance (SR 26-2), fair lending, fiduciary duties, and examination expectations. Ability to translate regulatory requirements into operational controls.
- **Technical credibility:** Sufficient understanding of AI and ML to evaluate vendor claims, assess validation approaches, and recognize when technical teams are deflecting or oversimplifying. This does not require the governance lead to be a data scientist, but they must be able to hold their own in technical discussions.

In practice, this dual competency usually requires a team rather than a single individual, or recruitment of rare individuals who bridge both domains. The team model pairs regulatory specialists with technical specialists under unified governance leadership.

Reporting Line Tradeoffs

Where the AI governance function reports determines its capture risk:

REPORTS TO	ADVANTAGE	CAPTURE RISK
CRO / Chief Risk Officer	Natural risk mandate, independence from IT	May lack technical access and credibility with engineering
CTO / CIO	Technical access, pipeline integration easier	Captured by delivery pressure; asked to approve what is already built
COO / Business Units	Business context, practical orientation	Captured by revenue pressure; governance becomes checkbox
General Counsel	Regulatory credibility, independence	May lack technical depth; legalistic rather than operational
CEO / Board Committee	Maximum independence and authority	May lack day-to-day operational connection; reserved for largest firms

The most effective structures typically have AI governance reporting to the CRO with a strong dotted line to the CTO for technical liaison, or a dedicated AI Governance function with direct Board committee access for Tier 1 escalations.

The “Effective Challenge” Test

The revised guidance retains “effective challenge,” critical analysis by objective, informed parties. For AI governance, apply this test to your organizational design:

- Can your governance function say “no” to a Tier 1 deployment the CTO wants to ship?
- Can your governance function obtain technical documentation IT does not want to provide?
- Does your governance function learn about AI initiatives before procurement, or after deployment?
- Can your governance function compel remediation of findings, or only recommend?
- Does your governance function have the technical vocabulary to challenge vendor claims?

If the answer to any of these is “no,” your governance structure has a design flaw that will manifest during examination or incident response.

CHAPTER 7

Policy Framework

Effective AI governance requires a comprehensive policy framework that provides clear guidance while remaining flexible enough to accommodate evolving AI capabilities.

Core Policies Required

1. AI Acceptable Use Policy

- Permitted uses by role.
- Prohibited uses (customer data in unapproved systems, regulated decisions without approval).
- Shadow AI prohibition.
- Vendor AI restrictions.

2. AI Development Policy

- Development lifecycle requirements by tier.
- Documentation standards (training data, architecture, hyperparameters).
- Testing and validation gates before deployment.
- Approval workflows and sign-off requirements.

3. AI Validation Policy

- Validation scope and frequency by tier.
- Validation techniques (conceptual soundness, performance, bias, explainability).
- Independence requirements.
- Findings classification and remediation timelines.

4. AI Monitoring Policy

- Continuous monitoring requirements (drift detection, performance metrics).
- Periodic review cadence by tier.
- Escalation triggers and thresholds.
- Incident response procedures.

5. Third-Party AI Policy

- Due diligence requirements before engagement.
- Minimum evidence package by tier.
- Ongoing oversight and performance monitoring.
- Concentration risk limits.

PART IV

Lifecycle Management

CHAPTER 8

Development and Validation

Pre-Development Requirements

Before AI development begins, several assessments should be completed:

- **Use case justification:** Clear articulation of business problem and why AI is appropriate.
- **Data availability assessment:** Confirmation that appropriate training data exists.
- **Regulatory impact analysis:** Identification of applicable regulations.
- **Ethical review:** For Tier 1 and 2 systems, assessment of potential bias and fairness concerns.

Development Documentation Standards

Model risk management guidance requires thorough documentation of model development. For AI systems, this must include:

- **Training data specification:** Source, selection criteria, preprocessing steps, known limitations.
- **Model architecture documentation:** Architecture selection rationale, alternatives considered.
- **Hyperparameter selection:** Parameters chosen and rationale for selection.
- **Feature engineering:** Features used, transformations applied, features excluded and why.
- **Version control:** Code versions, data versions, ability to reproduce training.

Validation Framework

Validation requirements vary by tier:

VALIDATION ACTIVITY	TIER 1	TIER 2	TIER 3	TIER 4
Conceptual soundness review	Required	Required	Recommended	N/A
Data quality assessment	Required	Required	Required	N/A
Performance testing	Required	Required	Basic	N/A
Bias and fairness testing	Required	Risk-based	N/A	N/A
Explainability assessment	Required	Risk-based	N/A	N/A
Independent validation	Required	Required	Self-assess	N/A

CHAPTER 9

Deployment and Monitoring

Deployment Controls

- **Staged rollout:** Phased deployment with monitoring gates before full production.
- **Rollback procedures:** Documented ability to revert to prior versions quickly.
- **Performance baseline:** Establish baseline metrics before deployment for comparison.
- **Human oversight configuration:** Configure appropriate human review for the tier level.

Ongoing Monitoring

REVIEW TYPE	TIER 1	TIER 2	TIER 3	TIER 4
Performance review	Monthly	Quarterly	Annual	N/A
Bias re -assessment	Quarterly	Annual	N/A	N/A
Full revalidation	Annual	18 months	As needed	N/A
Documentation refresh	Annual	Annual	As needed	N/A

Audit Trails: Logs versus Evidence

A critical distinction exists between operational logging and authorization artifacts:

CONCEPT	DEFINITION	REGULATORY VALUE
Logs	Operational trace of system activity: timestamps, user actions, errors, performance metrics	Necessary but not sufficient; can show “what happened” but not “why it was permitted”
Authorization Artifacts	Replayable, version-bound, policy-bound authorization records with tamper-evident integrity	Independently verifiable; can answer “how do you know?” for auditors, regulators, courts

For Tier 1 systems, the goal is to move beyond “audit logging enabled” toward tamper-evident authorization artifacts that show whether an effect-bearing action was permitted before release under the applicable policy, authority chain, provenance condition, and replay basis. This distinction becomes critical at maturity Level 5.

Observability Is Not Authorization

Logging, monitoring, and post-hoc supervision can show that an event occurred. They do not establish that the action was permitted before execution. For Tier 1 AI systems, the control objective is pre-execution authorization: the system must obtain an authorization verdict before an effect-bearing action is released. If the authorization artifact cannot be produced, the governed action fails closed.

Monitoring, model-performance tracking, telemetry, and visibility may support model-risk oversight, validation, audit, and examination readiness. Combining them may improve the evidence available to review, but it does not itself produce a determination made before release of what was permitted. For effect-bearing AI actions, require the authorization artifact, not only monitoring data or logs.

PART V

Audit and Examination Readiness

CHAPTER 10

Documentation Standards

Model Documentation Package (Tier 1)

Each Tier 1 AI system should have a complete documentation package containing:

- Model identification (unique ID, owner, developer, validator, tier rationale).
- Business context (use case, process integration, decision impact, regulatory touchpoints).
- Technical documentation (architecture, training data, features, hyperparameters).
- Validation documentation (scope, test results, findings, limitations, approvals).
- Monitoring documentation (approach, metrics, thresholds, escalation procedures).

Examiner Pack: Standard Binder Structure

For each Tier 1 model, maintain a ready-reference binder (physical or electronic) with these sections to minimize scrambling during examinations:

TAB	CONTENTS
1	Inventory entry and tier classification rationale
2	Use case description and regulatory touchpoints matrix
3	Validation summary with sign-offs (most recent and prior)
4	Monitoring dashboard, thresholds, and last 3 months evidence
5	Change log and version history
6	Issues and incidents, with CAPA (corrective and preventive actions)
7	Vendor artifacts (if applicable): model card, attestations, SLAs

CHAPTER 11

Examination Preparation

Common Examination Questions

Prepare responses to common examiner inquiries:

Inventory and Governance

- “Show me your complete AI and ML model inventory.”
- “How do you identify shadow AI?”
- “Walk me through your tiering methodology.”
- “Who has authority to approve Tier 1 models?”

Development and Validation

- “Show me validation documentation for [specific model].”
- “How do you test for bias in credit models?”
- “What is your approach to explainability?”
- “How do you validate third-party models?”

Monitoring and Control

- “Show me evidence of ongoing monitoring.”
- “What triggers a model review?”
- “How do you detect model drift?”
- “Walk me through a recent model issue and remediation.”

Examination Response Playbook

- **Documentation retrieval:** Know where all documentation lives and who can access it.
- **Subject matter experts:** Pre-identify SMEs for each major AI system.
- **Response coordination:** Designate an examination coordinator to manage requests.
- **Issue escalation:** Define an escalation path for unexpected questions or findings.

CHAPTER 12

Board and Executive Reporting

AI Governance KRIs and KPIs

Executives adopt what they can track. The following metrics support board-level AI governance reporting:

METRIC	DEFINITION	TARGET
Inventory coverage	Percent of known AI use cases tiered and documented	>95%
Tier 1 validation currency	Percent of Tier 1 models with current validation (under 12 months)	100%
Drift alerts	Count of material drift detections by severity (High/Med/Low)	Trending down
Bias exceptions	Count of bias and fairness testing exceptions requiring remediation	Zero for credit models
Vendor concentration	Percent of Tier 1 models dependent on single vendor	<30%

METRIC	DEFINITION	TARGET
Shadow AI detections	Count of unapproved AI use discovered per period	Trending down
Incident count	AI-related incidents by severity	Zero P1/P2
Time-to-remediate	Average days from finding to remediation by severity	<30 days for High

Report these metrics quarterly to the Model Risk Committee and annually (with trends) to the Board Risk Committee.

PART VI

Implementation Roadmap

CHAPTER 13

90-Day Implementation Plan

Phase 1: Foundation (Weeks 1 to 4)

Week 1 to 2: Inventory and Gap Assessment

- Conduct AI system inventory sweep across all business units.
- Map existing MRM framework to AI requirements.
- Identify policy gaps.
- Assess organizational readiness (skills, tools, governance).

Week 3 to 4: Governance Structure

- Expand MRM committee charter to include AI and ML.
- Define roles and responsibilities.
- Establish tiering criteria including automatic Tier 1 triggers.
- Draft core policies.

Phase 2: Critical Systems (Weeks 5 to 8)

Week 5 to 6: Tier 1 System Focus

- Identify and prioritize Tier 1 AI systems.
- Conduct gap analysis against SR 26-2 requirements.
- Develop remediation plans.
- Begin enhanced documentation.

Week 7 to 8: Validation Enhancement

- Implement bias testing for credit and fair lending models.
- Establish explainability requirements.
- Enhance monitoring capabilities.
- Compile vendor minimum evidence packages.

Phase 3: Operationalization (Weeks 9 to 12)

Week 9 to 10: Policy Deployment

- Finalize and approve policies.
- Conduct training for first and second line.

- Implement acceptable use controls.
- Establish shadow AI detection.

Week 11 to 12: Examination Readiness

- Compile Examiner Pack binders for Tier 1 systems.
- Conduct mock examination.
- Remediate gaps.
- Establish ongoing monitoring cadence and board reporting.

CHAPTER 14

Maternity Progression

Five Levels of AI Governance Maturity

LEVEL	NAME	WHAT IT LOOKS LIKE	TIMELINE
1	Foundational	Basic inventory exists, policies drafted but inconsistently applied	Current state
2	Structured	Complete inventory with tiering, policies approved, validation for Tier 1	6 to 12 months
3	Managed	Automated inventory, consistent validation, continuous monitoring, examination-ready	12 to 24 months
4	Optimized	Real-time risk visibility, automated compliance checking, predictive risk identification	24+ months
5	Independently Verifiable	The organization is designed to produce tamper-evident authorization artifacts for consequential AI actions, with evidence sufficient for independent replay	Infrastructure investment

Level 5: Independently Verifiable

At Level 5, the institution is designed to produce tamper-evident authorization artifacts for consequential AI actions, with evidence sufficient for independent replay. The maturity endpoint is not better monitoring. It is an independently reconstructable verdict: whether the action was authorized before it executed.

This is the logical endpoint of the authorization-boundary principle: governance is enforced at the point where an output becomes an effect-bearing action, and evidence, not explanations, must exist. Institutions pursuing Level 5 maturity should evaluate infrastructure capable of producing tamper-evident authorization artifacts, fail-closed release behavior, and replayable audit trails that support regulatory evidence needs.

TECHNICAL FOUNDATION

The Five Tests Standard

OPEN SPECIFICATION

The evidence standard in this guide aligns with the Five Tests Standard (5TS), an open specification for verifiable AI governance. 5TS defines five requirements: Stop, Ownership, Replay, Escalation, and Provenance. Provenance means origin, not truth. It is a test, not a verdict.

The verdict space is ALLOW, DENY, and ABSTAIN. An ABSTAIN blocks execution pending authorized human override; escalation is the consequence of ABSTAIN, not a fourth verdict.

Five Tests Standard (5TS) · Open standard DOI: [10.5281/zenodo.21040295](https://doi.org/10.5281/zenodo.21040295) · Explanatory foundation DOI: [10.5281/zenodo.21048661](https://doi.org/10.5281/zenodo.21048661)

APPENDICES

Templates and Reference

APPENDIX A

AI System Intake Form

Every AI use case must submit this form before deployment.

FIELD	DESCRIPTION / OPTIONS
Use Case Name	Free text, descriptive name
Business Owner	Name, title, business unit, single accountable person
AI System / Tool	Vendor or product name, or custom-built description
Governance Track	Decisioning AI / Communications AI / Both
Purpose / Outcome	What problem does this solve? What decisions or outputs?
Output Type	Decisioning / Drafting / Analytics / Communication
Data Classes	Public / Internal / Confidential / PII / PHI / MNPI
Regulatory Touchpoints	SR 26-2 / Fair Lending / Reg BI / 17a-4 / EU AI Act
Automatic Tier 1 Trigger?	Yes (specify which) / No
Proposed Tier	Tier 1 / Tier 2 / Tier 3 / Tier 4
Tier Justification	Brief explanation

APPENDIX B

Tier 1 Control Checklist

All controls required before Tier 1 AI system go-live:

- Designated business owner with documented accountability.
- Written purpose statement and expected business outcome.
- Automatic Tier 1 trigger assessment documented.
- Data classification completed and documented.
- Training data provenance documented.
- Authorization artifact evidence configured for effect-bearing actions, not only logging.
- Pre-execution authorization gate configured for all effect-bearing actions; human approval required where policy, regulation, or escalation rules designate.

- Incident response path defined.
- Model validation complete (conceptual soundness, performance, bias).
- Explainability approach documented (required for fair lending).
- Vendor minimum evidence package collected (if third-party).
- Legal and compliance sign-off on file.
- MRC approval documented.
- Examiner Pack binder compiled.
- Quarterly review scheduled.

APPENDIX C

Regulatory Reference Summary

REGULATION	SCOPE	KEY AI REQUIREMENTS
SR 26-2 / OCC 2026-13	Banks, BHCs (>\$30B)	Risk-based model governance, validation, ongoing monitoring, third-party model accountability; generative and agentic AI out of scope
UK SS1/23	UK-regulated banks and designated firms	Technology-neutral model-risk expectations; proportional implementation; applies across model and risk types, including in-house and vendor models
SEC Advisers Act	Investment advisers	Fiduciary duty, disclosure, enforcement against misleading AI claims
FINRA 3110 / 3120	Broker-dealers	Supervision of AI communications, recordkeeping
SEC 17a-4	Broker-dealers	Retention of AI-generated communications
ECOA / Reg B	Credit decisions	Fair lending, adverse action notices, explainability required
EU AI Act	EU market access / EU use cases	High-risk classification for creditworthiness evaluation of natural persons and life and health insurance risk assessment and pricing; conformity assessment, documentation, oversight, robustness, and database obligations where applicable

For questions on implementation or to discuss evidence-based governance infrastructure: [ferz.ai](#)