

Responsible AI Governance for Wealth Management

A practical operating framework for firms adopting AI inside real fiduciary obligations, without a model risk team.

WHO THIS IS FOR: CEOs, COOs, CCOs, CIOs/CTOs and senior leaders responsible for AI adoption, operations, risk and oversight in RIAs and family offices.

The paradox at the center of this framework

AI can amplify human capability while simultaneously weakening the judgment needed to supervise it.

Matt Chambers

718 Innovations LLC

Practitioner Framework | August 2026

Practitioner guidance only. This framework is not legal, regulatory, or compliance advice.

© 2026 718 Innovations LLC.

Executive Summary

AI is entering wealth management firms faster than most governance structures are adapting. It arrives through employee use, vendor platforms, embedded features, and increasingly autonomous technology. The governance challenge is not whether a firm should use AI, but whether leadership knows where it is being used, what it touches, what consequences can follow, and who is accountable.

The timing is not theoretical. The SEC's 2026 Examination Priorities explicitly include automated investment tools and AI, reinforcing the need for firms to understand where AI is being used and how it is governed.

Responsible AI governance should enable adoption, not obstruct it. The objective is enough visibility, accountability, and proportionate control for lower-risk uses to move quickly while applying greater scrutiny where AI touches sensitive information, fiduciary judgment, material client outcomes, or actions that are difficult to reverse.

Six practitioner conclusions

1. **Start with visibility, not policy.** A firm cannot govern AI use it cannot see. Establish what is being used, for what purpose, with what data, and with what consequences.
2. **AI extends capability. It does not replace judgment or accountability.** AI can produce convincing answers without understanding consequences. Human oversight only works when the reviewer has enough applicable knowledge, context, and authority to challenge it.
3. **AI risk does not begin or end with the model.** Data quality, historical decisions, business processes, employee behavior, vendor design, and implementation can shape outcomes as much as the model itself.
4. **Govern according to consequence, not technology.** Oversight should increase with data sensitivity, impact, autonomy, explainability, reversibility, and fiduciary or regulatory significance.
5. **Technology can be outsourced. Accountability cannot.** Every material use needs clear ownership and decision rights across leadership, operations, compliance, technology, and the business.
6. **Build the minimum controls necessary to make informed decisions.** Governance should create guardrails rather than roadblocks. Excessive friction drives workarounds, shadow AI, and reduced visibility.

The first practical step

Establish visibility. Identify the highest consequence uses. Set immediate data boundaries. Assign accountability. Then evolve the controls with evidence.

Introduction

This framework assumes no dedicated AI team, model risk function, or large technology budget. It is designed to be utilized within the leadership, compliance, technology, and operating structures a firm already has. It builds on CFA Institute, NIST, SEC, and FINRA guidance and adds the practitioner question those sources do not answer: how can a firm without specialist resources practice good AI governance day to day?

Three AI categories used in this framework

1. Generative AI Creates or transforms content: text, summaries, meeting notes, images, and communications.	2. Predictive and analytical AI Identifies patterns, produces forecasts, prioritizes opportunities, or supports recommendations from available data.
3. Agentic AI Goes beyond producing information and can initiate or complete actions with varying levels of human involvement.	WHY THE DISTINCTION MATTERS These categories describe what the technology does, not how risky it is. Risk depends on what the system can access, influence or do, and what follows if it is wrong. A tool that summarizes an internal meeting does not create the same exposure as one that recommends a client action, or one that takes it.

Central premise

The purpose of AI governance is to enable responsible adoption, not prevent it. It separates where AI can move quickly from where consequence requires scrutiny, accountability, and control.

The Operating Problem

Artificial intelligence is moving into the everyday operations of wealth management firms through productivity tools, CRM platforms, marketing applications, analytics, client service workflows, third-party technology, and employee use of public AI tools.

For leaders of RIAs, family offices, and other regulated financial services businesses, the question is no longer whether AI will enter the firm. It already has. The question is whether anyone can describe where it is, what it touches, and who answers for it.

A firm that cannot answer all six has a governance gap. That gap, rather than the technology itself, is the problem this framework addresses.

1. Where is AI being used?	2. What business problem is it intended to solve?
3. What data and decisions does it touch?	4. What could go wrong?
5. How much oversight is appropriate?	6. Who is accountable for the outcome?

Start with visibility

The first step in responsible AI governance is not writing a policy. It is knowing where AI is already being used, what it has been asked to do, what information it can reach, and what follows if it is wrong.

1. AI Is Already Inside the Firm

AI adoption rarely begins with an enterprise strategy. It arrives incrementally: an employee improves a client email, marketing adopts AI-assisted content, a CRM vendor adds summarization, operations automates part of a workflow, or a provider enables an AI feature during a routine upgrade.

Each decision looks small. Together they create an AI environment leadership may not fully understand or be able to describe.

IN PRACTICE

Leadership confidence about AI use often masks two opposing blind spots:

- 1. **Informal AI adoption is underestimated.**
- 2. **The sophistication of technology already in use is overstated.**

Industry pressure makes both worse: leaders are expected to show AI progress, which encourages both hidden experimentation and overstated sophistication.

The danger is governing a version of AI use that does not match reality. Approved tools can mask shadow use, while conventional automation labeled as AI can overstate capability.

The first task is visibility: what is in use, where, and with what consequences.

2. Understanding What AI Is Not

Most discussion of AI concerns what it can do. Firms also need to understand what it is not. That is what allows leaders to separate capability from hype, confidence from understanding, and useful assistance from decisions that still require human judgment.

AI is not a person

AI communicates in ways that can appear thoughtful, confident, and at times empathetic. It summarizes complex information, identifies patterns, and produces recommendations that resemble human analysis. But it does not possess judgment, intent, experience, discernment, or professional accountability.

The risk is that users begin to treat AI as though it possesses those qualities. Polished, confident output can create an impression of authority the system does not hold. Users may accept an answer because it sounds informed rather than because they have tested its reasoning, evidence, or applicability.

Perceived knowledge is not applicable knowledge

Immediate access to an answer is not the same as understanding the subject well enough to evaluate it or act on it responsibly.

That distinction matters throughout this framework because bias, handoffs, and weak review can all produce apparent knowledge without matching applicable knowledge.

AI is not inherently accurate

Generative AI can produce convincing answers that are incomplete, misleading, or wrong. Predictive systems can reproduce weaknesses in historical data and business practice. Agentic systems can execute incorrect actions quickly and at scale. Risk rises when reviewers lack the knowledge or context to recognize the error.

The hardest errors are often plausible ones. They are coherent, consistent with expectations, and therefore less likely to be challenged. Bias in the data, process, system, or human review can make them harder to recognize.

A core distinction

Plausibility is not accuracy. Access to information is not knowledge.

AI does not understand consequence

AI processes information according to its design, data, instructions, and available context. It does not experience the operational, financial, regulatory, or human consequences of a recommendation or

action. In wealth management, a small error can reach a client's financial position, privacy, tax planning, investment decisions, or trust in the firm.

AI can generate a recommendation in seconds. Understanding what happens when you implement it can take years to learn.

Automation does not transfer accountability

Greater autonomy reduces the human effort a task requires. It does not reduce the firm's responsibility for the outcome. AI can extend capability, support judgment, and automate activity, but it cannot inherit accountability or supply the applicable knowledge needed to exercise that accountability well.

The people who remain in the loop need enough applicable knowledge, discernment, and authority to challenge the system. Leaders need enough understanding to judge where AI creates value, where it can mislead, and where human judgment must remain decisive.

IN PRACTICE

A recommendation can be logical on screen and still be wrong in practice. An AI system may recommend reducing a concentrated position while missing tax consequences, estate-planning intent, transfer restrictions, family dynamics, or information never captured in the data it received.

Experienced professionals notice not only what is present in the analysis, but what is missing. That discernment comes from applicable knowledge, context, and experience the system may never have been given.

The recommendation is the easy part. Understanding the consequence is the work.

Two misconceptions repeatedly distort how firms assess AI:

1. **Automation is mistaken for AI.**
Rules-based workflows, scoring, decision trees, and automated alerts can create substantial value, but they are not modern AI simply because they are automated or marketed that way. Modern AI derives part of its capability from patterns learned from data rather than solely from rules programmed in advance.
2. **A convincing or authoritative answer is mistaken for a correct one.**
Generative AI allows vendors and employees to produce polished answers without equivalent depth or applicable knowledge. Confidence is not evidence of correctness or understanding. Depth becomes clear when assumptions, reasoning, limitations, and consequences are tested.

Both create the same problem: perceived capability and perceived knowledge exceed what actually exists.

The risk is that leadership may then make decisions with a level of confidence that the underlying capability or expertise does not justify.

3. AI Risk Does Not Begin With the Model

A well-designed system can still produce poor outcomes if it draws on weak data, reflects a flawed business process, is implemented badly, or is used by people who cannot challenge it. The model is only one input among several.

Seven sources of bias

Each of the following can shape what an AI system produces and what users afterward treat as reliable knowledge.

1. Data quality Incomplete, inaccurate, inconsistent, or outdated data distorts outputs before analysis begins. AI processes poor data faster. It does not make poor data reliable.	2. Historical decisions AI can reproduce outdated assumptions, inconsistent practices, and prior human bias embedded in past decisions. History describes what happened. It does not establish what should happen next.
3. Business process Bias lives in workflows, rules, incentives, handoffs, and default settings. AI can reinforce an existing process weakness rather than correct it.	4. Adviser and employee People shape AI through the questions they ask, the information they supply, the outputs they accept, and what they do next. Oversight only works when the reviewer can meaningfully challenge the output.
5. Demographic AI can produce different outcomes across groups even where no such difference was designed in. Test whether differences are explainable and justified.	6. Vendor model Third party AI introduces a distinct source of risk. Vendor assumptions, model design, training data, and updates may not be visible. Using someone else's system does not transfer accountability.
7. Implementation Configuration, thresholds, defaults, integrations, permissions, training, and local workarounds change how a sound model behaves in practice.	THE PATTERN WORTH NOTICING Only one of these seven is explicitly about the model. The firm has greater ability to act on the environment around the model than on the model itself.

Bias can distort both the answer and the discernment required to challenge it.

Risk compounds through interpretive handoffs

Business intent rarely moves directly from a subject matter expert or senior leader into an AI system. It travels through people and processes.

Business intent → requirements → system configuration → documentation → training → daily practice
Meaning shifts at every handoff. No single step has to fail dramatically for the end result to differ materially from the original intent. Small losses of context accumulate, and everyone along the chain believes they represented it faithfully.

The apparent intelligence of the finished system can conceal weaknesses introduced earlier in the chain. The same effect operates on organizational knowledge: as people generate answers without building the knowledge required to evaluate them, activity gets faster while discernment gets thinner.

IN PRACTICE

Loss of meaning between business intent and technology is not new. Traditional software projects have always translated objectives into requirements, rules, configurations, and user behavior.

AI amplifies the problem because more meaning is inferred and less of the decision path is explicitly visible. Each step can appear reasonable while the connection between original intent and resulting output becomes harder to trace.

The danger is that AI can conceal the loss behind output that still looks intelligent and convincing.

The objective is not to eliminate every source of bias or error. It is to know where risk enters, how it compounds, and where enough knowledge, oversight, and control are needed before risk becomes consequence.

4. Not All AI Deserves the Same Governance

Effective governance is proportionate to consequence. Treating every AI use as equally risky creates friction around routine work while leaving too little attention for uses that can materially affect clients, employees, or the firm.

Too much control produces less control

Bureaucratic, intrusive, or slow governance creates pressure to route around it. The consequences include shadow AI, undocumented workarounds, delayed consultation, reduced visibility, and weaker accountability. The aim is guardrails, not roadblocks.

Governance should rest on clear guidelines, defined boundaries, accountable ownership, and proportionate oversight rather than approval at every step. Some boundaries remain non-negotiable, particularly where AI touches sensitive client data, fiduciary decisions, regulated communications, material client outcomes, or actions that are difficult to reverse.

Classify by consequence, not by technology

Assess each material AI use against six practical factors: data sensitivity, impact, autonomy, explainability, reversibility, and fiduciary or regulatory significance. These matter more than what the technology is called.

Risk tier	Typical uses	Governance approach
Low	Internal drafting, summarization, research, or administrative assistance where outputs are reviewed and errors are readily reversible.	Approved tools, employee training, clear guidelines, routine oversight within established boundaries.
Moderate	Client communications, segmentation, workflow prioritization, or analytical recommendations that may influence service or business decisions.	Named owner, documented testing, defined human oversight, appropriate compliance involvement, periodic monitoring.
High	Sensitive client information, material client outcomes, significant automated actions, or systems acting across multiple applications.	Executive and compliance approval, formal validation, enhanced documentation and controls, monitoring thresholds, explicit stop conditions.

The same technology can carry very different risk

An AI system identifies clients who might benefit from considering a Roth conversion. If it surfaces the opportunity to a qualified adviser who independently evaluates the client's circumstances, the use is moderate risk. The same technology becomes something else entirely if it determines the conversion is appropriate, communicates individualized advice to the client, modifies a financial plan, or initiates an action without meaningful professional review.

The technology is identical. The consequence, the autonomy, and the role of human judgment are not.

Human involvement has to be meaningful

Requiring a person to approve every AI output does not create oversight. Review still fails when the reviewer introduces bias, accepts output without scrutiny, or simply gets used to clicking approve. Oversight becomes meaningful only when the reviewer has enough applicable knowledge, information, time, and authority to challenge what is in front of them.

Validated systems can perform bounded activities with greater autonomy where performance is measurable, monitored, and reversible. Human judgment should become more decisive as decisions involve fiduciary judgment, significant financial consequence, sensitive information, or actions that are hard to undo.

Classifications have to evolve

Risk designations are not permanent. Changes to the model, vendor, data, integrations, autonomy, purpose, or operating environment can alter risk. Reassess after material change and on a regular cadence; strengthen controls where experience shows greater risk and simplify those that add little.

5. Accountability Stays With the Firm

AI changes how work is performed, not who answers for the outcome. Providers may supply the model, infrastructure, or expertise; the firm remains responsible for how those capabilities are used inside its business.

Technology can be outsourced. Accountability cannot.

Effective governance needs named decision rights, not broad statements that the business, compliance, or technology owns AI.

Role	What they own
CEO and firm leadership	Risk appetite, strategic priorities, and boundaries for AI adoption.
COO or senior operating leader	Cross functional ownership connecting business objectives, process, technology, people, and risk.
CCO or compliance leader	Where fiduciary, regulatory, supervisory, privacy, communication, and recordkeeping requirements apply.
Technology and information security	Technical standards for systems, integrations, access, security, and data protection.
Business owner	The intended outcome, success measures, material risks, and intervention conditions.
Employees and advisers	Use of approved tools within guidelines, and professional judgment where the role requires it.

In a smaller firm, one person may hold several roles. Structure is not the issue. Responsibilities must be explicit, and accountable people need real authority to request testing, require remediation, escalate, modify a use, or suspend higher-risk activity.

Vendors can provide evidence about security, limitations, testing, and performance. They cannot determine whether a capability fits this firm's clients, data, workflows, and obligations; that decision remains with the firm.

IN PRACTICE

Ambiguous decision rights create meetings and approvals without clear ownership. Effective governance requires the right people to hold the knowledge, authority, and accountability to decide, not everyone to approve everything.

6. The Minimum Required Operating Controls

An end state, not a day one checklist

Build toward these in proportion to firm size, AI footprint, and risk; start with visibility, immediate data boundaries, higher-risk activity, and clear accountability.

These six disciplines form a practical foundation for responsible AI use.

Discipline	What it requires
1. Accountability and approval	Each material use has a named owner accountable for the intended outcome. Low risk uses run within guidelines. Review increases with consequence, autonomy, and regulatory significance. The point is not to collect approvals. It is that someone knowingly accepted the risk.
2. Data and privacy boundaries	Employees can tell approved integrated AI, approved enterprise AI, and public or unapproved AI apart. Client NPI, credentials, and proprietary firm data stay out of environments the firm does not control for storage, processing, retention, or reuse.
3. Testing and meaningful oversight	Testing is proportionate to the consequence of failure and may cover accuracy, completeness, consistency, bias, security, privacy, integrations, explainability, reversibility, and behavior under stress. External specialists can help. The decision stays with the firm.
4. Transparency and traceability	For material AI supported decisions the firm can establish what system was used, what information it relied on, what role AI played, what human review occurred, what action followed, and who was accountable. The objective is decision traceability, not paperwork.
5. Vendor and technology governance	Diligence goes past cost, functionality, and cybersecurity to what the capability does, what data it uses and retains, whether firm or client data trains the model, how performance is tested, how changes are communicated, and how incidents are handled.
6. Monitoring, escalation, stop conditions	Deployment is not the end. Monitor errors, overrides, complaints, unusual outcomes, incidents, vendor changes, and performance against the business objective. Higher risk uses carry explicit conditions for investigation, restriction, remediation, or suspension.

Controls must work inside the way work actually happens. A control that consistently obstructs legitimate work creates pressure to bypass it; a bypassed control can be worse than none because it creates the appearance of oversight without the substance.

Minimum necessary control

Use the least control needed to support an informed decision, maintain accountability, and protect against consequences the firm is not willing to accept.

7. Implementing Without Creating Bureaucracy

A firm does not need the complete end-state framework before it can use AI responsibly. Attempting every control and review process at once creates burden and slows the progress the framework exists to support.

The hardest part is getting started

The subject looks larger and more technical than it needs to be. Do not design the perfect governance model first; build enough visibility and structure to take one informed step.

Phase	What happens
1. Visibility and immediate boundaries	Identify known AI uses, surface obvious shadow use, set clear restrictions around sensitive data, designate a governance lead, flag anything needing immediate review.
2. Classify and prioritize	Classify material uses against the risk factors. Let lower risk uses move within guidelines. Prioritize moderate and high risk activity. Ask whether a use creates enough value to justify its complexity.
3. Build the required controls	Define ownership, data boundaries, testing, human oversight, vendor requirements, monitoring, and escalation for the risks that actually exist. Highest consequence first.
4. Expand responsibly	Scale successful lower risk uses. Advance moderate risk as evidence supports it. Move higher risk forward only with enough testing, understanding, ownership, and oversight to decide well.
5. Establish cadence	Review new uses, changed classifications, vendor changes, incidents, overrides, complaints, bypassed controls, and performance against outcomes. Fit this into the firm's existing operating rhythm rather than building a separate management system for AI.

From framework to how we do things

As the framework matures, responsible AI use becomes part of how the firm operates. Employees know which tools and data are appropriate, when they can act within guidelines, when they need help, and when consequence requires escalation.

Repeated behavior becomes habit and habit becomes culture. Routine decisions require less intervention, freeing leadership to focus on the new, unusual, and consequential. Strong governance eventually becomes less visible because informed decision making has become ordinary.

Deliberate governance also creates evidence

A deliberate governance framework leaves the firm able to show not only what decisions it made, but how and why. SEC rules and examination priorities emphasize reasonably designed policies and procedures, effective implementation, governance, controls, third-party oversight, and emerging

technology risk. The framework therefore creates evidence of informed decision making, accountability, oversight, and evolution for regulatory examination and sophisticated client diligence.

The standard is not perfection

The objective is not to show that every decision was right. It is to show that important decisions were informed, deliberate, accountable, and appropriately reviewed.

The cost of not governing

Governance consumes capacity, but so does unmanaged AI: duplicated subscriptions, shadow use, unknown data exposure, poor implementation, underused vendor capability, operational errors, remediation, and slower adoption because leadership lacks confidence about what is safe.

A useful test

The question is not whether governance has a cost. It is whether that cost buys enough visibility, confidence, and control to enable responsible value creation.

Closing Perspective: Start Deliberately

AI adoption is outpacing governance. Schwab's 2026 research found 63% of advisers already using AI tools and a further 23% considering them, while only about one in five said their firm had a vision for AI adoption.

63%	23%	about 1 in 5
already using AI	considering AI	have an AI vision

Source: Schwab Advisor AI in Action research, 2026. Survey of 533 advisers who custody at Schwab.

For a firm without a deliberate approach, waiting is not neutral. AI keeps arriving through employees, vendors, and embedded systems while leadership remains unclear on use, data, consequence, and accountability. The practical response is to take the first informed step, not to complete every element of this framework before acting.

Get started deliberately

Establish visibility. Identify the highest consequence uses. Set immediate boundaries. Assign accountability. Then evolve with evidence.

Build capability, not dependence

Outside expertise can accelerate initial implementation, evaluate higher-risk uses, or support specialized testing. Where internal capacity is thin, external help can establish structure faster and avoid

predictable mistakes. The objective should be to transfer practical capability into the firm, not create permanent dependence.

If this is useful to you

I work with firms through a short diagnostic to establish visibility and immediate boundaries; hands-on build of the operating model and controls; and periodic review as technology, vendors, and regulation change.

If this framework is useful, I am glad to talk it through.

Matt Chambers, 718 Innovations LLC.

About the Author

Matt Chambers is a wealth management operating executive and founder of 718 Innovations LLC, which he formed in 2020 as an advisory vehicle for select independent engagements. His career has remained centered on executive operating leadership including a \$17 billion multi-family office and founder-led advisory firms, with responsibility spanning client and investment operations, finance, people, technology, risk, growth, and enterprise change.

He has served as COO, Chief People Officer, and CIO, partnering with founders, CEOs, and senior leaders to translate strategy into accountable execution and build the operating capability required to scale. Earlier in his career, he helped grow a financial-services consulting practice from six people to eighty and led complex transformation programs across institutional financial services.

His applied AI work in regulated financial services dates to 2017 and spans predictive, generative, and agentic use cases, including workflow design, implementation, and the governance required for responsible adoption. His formal study includes Strategic AI at Cornell University and an MBA concentration in Artificial Intelligence at Rome Business School.

This practitioner framework draws on that applied experience, formal study, current regulatory guidance, and external evidence.

Selected References

- Preece, R. (2022). Ethics and Artificial Intelligence in Investment Management: A Framework for Professionals. CFA Institute Research and Policy Center.
- Tabassi, E. (2023). Artificial Intelligence Risk Management Framework (AI RMF 1.0), NIST AI 100-1. National Institute of Standards and Technology.
- National Institute of Standards and Technology (2024). AI Risk Management Framework: Generative Artificial Intelligence Profile, NIST AI 600-1.
- U.S. Securities and Exchange Commission, Division of Examinations (2025). 2026 Examination Priorities.
- U.S. Securities and Exchange Commission (2003). Compliance Programs of Investment Companies and Investment Advisers, Release Nos. IA-2204; IC-26299.
- U.S. Securities and Exchange Commission (2024). Regulation S-P: Privacy of Consumer Financial Information and Safeguarding Customer Information, Release No. IA-6604.

- FINRA (2024). Regulatory Notice 24-09: Regulatory Obligations When Using Generative Artificial Intelligence and Large Language Models.
- FINRA (2026). GenAI: Continuing and Emerging Trends, 2026 Annual Regulatory Oversight Report.
- Charles Schwab & Co., Inc. (2026). AI for RIAs: From Experiment to Strategy.

© 2026 718 Innovations LLC.