

AI Suitability without Black Boxes: A Governance Model for Relationship Managers

An International Evidence, Decision and Oversight Framework for AI-Assisted Wealth Recommendations

Authored by Chennakeshav Adya

Independent Researcher

August 2026

Abstract

Artificial intelligence can help relationship managers organise client information, retrieve product evidence, compare mandates, identify inconsistencies and draft explanations. These uses can improve speed and consistency across a large product shelf. They can also create unsuitable recommendations at scale when client data are incomplete, product documents are stale, prompts are poorly controlled, models generate unsupported statements or staff treat an output as authoritative.

This paper develops an international governance model for AI-assisted suitability in private banking, external asset management and independent wealth advice. The model separates six elements of every recommendation: client evidence, product evidence, applicable policy, machine transformation, human judgement and client communication. It requires an evidence lineage that a qualified reviewer can reconstruct without access to proprietary model weights or source code. It also defines relationship-manager decision rights, validation, product governance, override controls, disclosures, record keeping, third-party oversight, monitoring, incidents and board reporting.

The framework draws on official materials from the United Kingdom Financial Conduct Authority, the European Securities and Markets Authority, the Hong Kong Securities and Futures Commission, the Australian Securities and Investments Commission, the United States Financial Industry Regulatory Authority, the Dubai Financial Services Authority, Singapore's Personal Data Protection Commission and the International Organization of Securities Commissions. The sources differ in legal status, scope and jurisdiction. They are used as design evidence and do not imply that every requirement applies to every wealth firm.

Six figures present the suitability decision stack, evidence-lineage chain, human control gates, model-risk heatmap, explanation card and monitoring dashboard. Six tables provide an evidence record, international comparison, decision-rights matrix, validation plan, hypothetical recommendation comparison and implementation roadmap. All scores, thresholds, client details, product attributes and scenarios in worked examples are hypothetical management assumptions created to demonstrate the method. They are not benchmarks, forecasts or legal conclusions. Actual obligations and decisions depend on jurisdiction, licence, client classification, service, product, contract and facts. This paper provides general information for professional audiences and does not provide legal, regulatory, technology, investment, tax or accounting advice.

JEL Classification: G23, G28, G41, L86, M15

Keywords: artificial intelligence, suitability, relationship managers, wealth management, investment advice, explainability, human oversight, model risk, client outcomes

1. Define what the machine is allowed to do

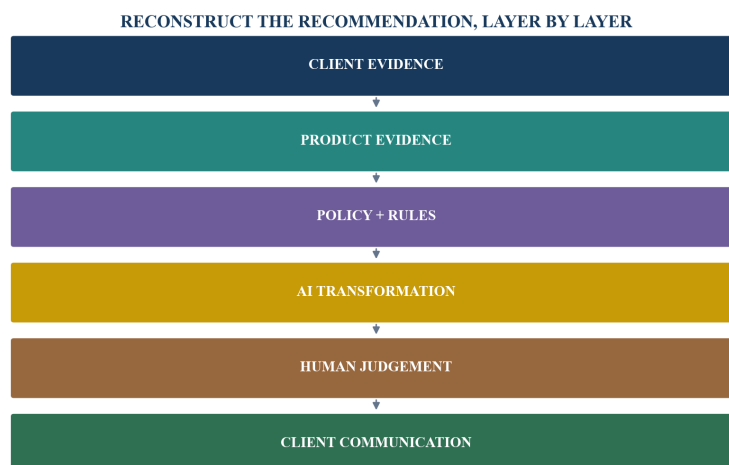
AI-assisted suitability should begin with a precise use case. A relationship manager may use technology to retrieve a client's stated objectives, consolidate holdings, extract product terms, compare documented constraints, surface inconsistencies or draft a plain-language explanation. Each task has a different risk. Retrieval can return the wrong document. Extraction can omit a condition. Comparison can apply the wrong rule. Generation can create a plausible statement with no source.

The firm should state whether the system informs, recommends, decides or communicates. It should identify the user, client population, products, data, model, interfaces, output and prohibited uses. A tool approved to summarise internal research should not automatically be used to recommend a complex structured product. A model validated for English documents should not be assumed to perform equally in another language.

Suitability remains an end-to-end professional process. The relationship manager must understand the client's circumstances and objectives, the product and its risks, the reason for the proposed match and the quality of the supporting evidence. Technology can augment that work. It cannot turn incomplete facts into a reliable recommendation.

The operating boundary should be expressed through permitted actions, required approvals and stop conditions. Missing source documents, stale client information, conflicting objectives, an unavailable risk disclosure or a product outside the approved shelf should stop automated progression. A control that merely displays a warning while allowing completion can be ineffective when staff face time pressure.

Figure 1. AI-assisted suitability decision stack



Every output is linked to sources, controls, human decisions and the client record.

Author framework. Each layer preserves its evidence, owner and decision status.

Table 1. Minimum suitability evidence record

Evidence class	Required content	Status field	Reviewer question
client	objectives, horizon, capacity, knowledge, experience, liquidity, restrictions and preferences	verified, client-stated, calculated, stale or missing	which facts support the proposed outcome?
portfolio	holdings, exposures, concentrations, cash flows, tax or legal constraints where relevant	current, reconciled or exception	how does the recommendation change total exposure?
product	terms, risks, costs, liquidity, scenarios, target market and conflicts	approved, current, superseded or restricted	which authoritative document supports each attribute?
policy	suitability rule, product-governance rule, delegated authority and required disclosure	applicable, conditional or unresolved	which rule permits, restricts or stops the action?
machine	model, version, prompt, input set, retrieval sources, output and confidence indicator	accepted, challenged, overridden or rejected	what did technology contribute and where can it fail?
human	analysis, alternatives considered, judgement, override, approval and communication	draft, approved, escalated or declined	could the approver explain and defend the result?

The record structure should be tailored to the applicable legal and conduct framework.

2. Treat explainability as evidence reconstruction

The phrase black box can refer to several different problems. A model may be mathematically complex. A vendor may protect its intellectual property. A relationship manager may receive a ranking without source citations. A compliance reviewer may see the final answer but lack the inputs and decision path. Governance should identify the actual opacity and address its effect.

For suitability, practical explainability means that a qualified reviewer can reconstruct why a product was considered, why alternatives were excluded, which client facts mattered, which product attributes mattered, what the system produced, what the human changed and what was communicated. This reconstruction does not require public disclosure of model weights. It requires decision-relevant evidence.

The record should preserve source identity, effective date, version, extraction time and transformation. If an AI assistant states that a product has monthly liquidity, the recommendation file should link to the approved document and provision supporting that statement. If it estimates portfolio exposure, the record should preserve positions, prices, method and timestamp. If it interprets a client's free-text objective, the RM should confirm the interpretation.

Explanations should match the audience. The client needs clear reasons, risks, costs and alternatives. The relationship manager needs actionable evidence and uncertainty. Compliance needs traceability and rule application. Model-risk staff need data, testing and failure analysis. Senior management needs outcome, incident and control trends.

3. Build evidence lineage from source to communication

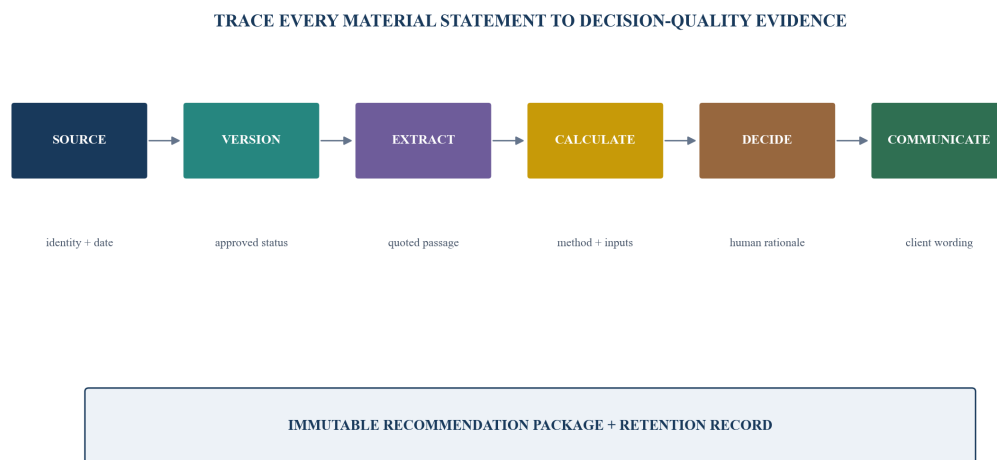
Evidence lineage tracks information from origin to the recommendation and final communication. It should identify client records, product documents, calculations, policies, model calls, retrieved passages, edits and approvals. A stable identifier for each item allows the firm to reproduce the record after source documents or models change.

Lineage controls should address freshness. Client circumstances can change. Product terms, costs, liquidity and risk classifications can be amended. Investment views and market conditions can move. The process should define how current each evidence class must be and what occurs when a source is outside that window.

Lineage should also preserve absence and conflict. A missing liquidity need must remain visible. Conflicting client answers should not be resolved silently by a model. The RM should seek clarification, record the resolution or stop the recommendation. Unknown information is a decision fact.

The system should separate source content from generated content. Retrieval passages, calculated fields and drafted narrative can be displayed differently. An assertion without an attributable source should be flagged for verification or removed before communication.

Figure 2. Evidence-lineage chain



Author framework. The evidence package preserves source, transformation, decision and communication records.

4. Put suitability obligations before model design

Official approaches differ by market. The FCA applies existing frameworks including the Consumer Duty and senior-management accountability to AI use. ESMA has stated that firms using AI in investment services should continue to comply with MiFID II obligations and has highlighted bias, opaque decision-making, overreliance, privacy and security risks. The Hong Kong SFC identifies investment recommendations, advice and research produced by generative AI language models as high-risk uses requiring additional controls.

FINRA reminds member firms that technology-neutral securities rules continue to apply to generative AI. ASIC's review of financial-services licensees identified a risk that governance may lag adoption, including gaps concerning fairness, bias and third-party systems. IOSCO's AI and machine-learning report addresses governance, testing, data quality, outsourcing, disclosure and human oversight for market intermediaries and asset managers.

Singapore's Model AI Governance Framework provides a voluntary, cross-sector design reference based on explainability, transparency, fairness and human-centric decision-making. DFSA materials and current supervisory communications provide evidence on governance, accountability, oversight and explainability within financial services. Each source has its own legal status and scope.

A firm should map the exact requirements applying to its entity, licence, service, client and product. The common design disciplines can support a group control model. Local legal analysis must determine suitability tests, disclosures, approvals, records, complaints, data protection and supervisory notifications.

Table 2. International design evidence for AI-assisted suitability

Market	Official material	Design implication	Important boundary
United Kingdom	FCA AI approach and Consumer Duty materials	preserve customer-needs, value, understanding and support outcomes; maintain senior accountability	existing rules apply according to firm, service and product scope
European Union	ESMA statement on AI in investment services	control bias, opaque outputs, overreliance, privacy and security while meeting MiFID II duties	EU AI Act and sector rules require separate applicability analysis
Hong Kong	SFC circular on generative AI language models	classify advice and recommendations as high-risk; apply lifecycle governance and additional safeguards	circular applies to licensed corporations within stated scope
Australia	ASIC Report 798	align governance with the scale and complexity of AI; address fairness, bias and third-party risk	findings reflect the reviewed licensees and existing obligations
United States	FINRA Regulatory Notice 24-09	technology use does not displace supervision, communications, records and conduct duties	notice addresses FINRA member firms and creates no new requirements
Singapore	PDPC Model AI Governance Framework	define roles, human involvement, operations controls and stakeholder communication	voluntary and cross-sector; sector requirements remain separate
DIFC	DFSA AI survey and supervisory materials	establish clear accountability, oversight, ethical data use and risk management	exact obligations depend on the DFSA Rulebook and authorised activity
International	IOSCO AI and machine-learning final report	govern development, testing, monitoring, data, providers, disclosure and human intervention	recommendations operate through member jurisdictions

5. Allocate decision rights to real people

Human oversight should be more than a click. The approver needs adequate knowledge, time, evidence and authority to challenge the output. The operating model should identify which decisions belong to the RM, product owner, compliance, investment committee, model owner, technology function and senior management.

The RM should own the accuracy of captured client information, the relevance of stated objectives, the explanation of alternatives and the professional rationale for the recommendation. The product owner should own approved product evidence, target market, restrictions and lifecycle events. Compliance should own rule interpretation, surveillance and escalation design. Model-risk or a suitably independent control function should own validation standards and material limitations.

Segregation matters. The same commercial team should not approve the use case, select validation evidence, judge performance and close its own deficiencies without independent challenge. Smaller firms can apply proportionality through committees or external specialists while preserving accountable decisions.

Delegated authority should be explicit. A low-risk, non-complex allocation within a documented mandate may follow a standard workflow. A concentrated, illiquid, leveraged or novel product may require additional product and compliance approval. A material system limitation may require temporary suspension of the AI function.

Table 3. Decision-rights matrix

Decision	Proposes	Challenges	Approves	Required evidence
approve an AI suitability use case	business and model owner	compliance, risk, information security and data protection	accountable executive or committee	use-case scope, risk assessment, validation and operating controls
add a product to the AI-enabled shelf	product owner	investment, compliance and operations	product committee	current terms, target market, risks, costs, liquidity and data mapping
make a client recommendation	relationship manager	workflow controls and second line where triggered	authorised RM or delegated approver	complete evidence record, alternatives and rationale
override a system warning	relationship manager	compliance or designated reviewer	authority set by severity	warning, reason, evidence, client impact and approval
change model, prompt or retrieval source	model or technology owner	validation, compliance and security	change authority	impact assessment, regression tests, version and rollback
suspend the system	incident lead, compliance or model owner	business owner where time permits	named incident authority	breach threshold, affected cases, containment and communication plan

Titles are illustrative; firms should allocate named accountable people under their governance framework.

6. Create meaningful human control gates

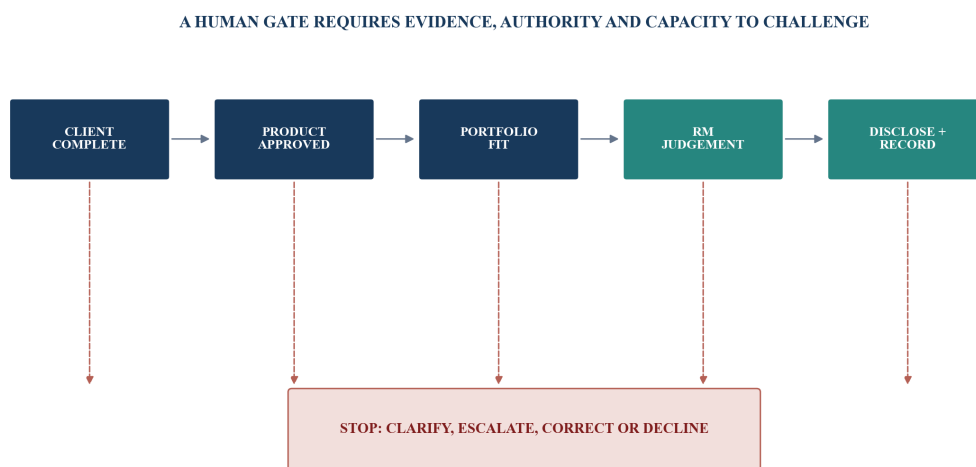
Control gates should occur before a recommendation becomes client-facing or executable. A first gate can test whether client evidence is complete and current. A second can confirm that the product is approved and evidence is current. A third can test portfolio fit, concentration, liquidity and costs. A fourth can require an RM to compare alternatives and record judgement. A final gate can confirm disclosure, approval and record retention.

Automation should stop when a material condition fails. Examples include inconsistent client objectives, missing capacity-for-loss information, an expired product document, unavailable cost data, an unapproved product, a concentration breach, an unsupported output or a system incident. The stop message should explain what evidence is missing and who can resolve it.

An override should be rare, specific and reviewable. It should identify the warning, decision maker, authority, supporting evidence, client impact and expiry. Repeated overrides may indicate that a rule is poorly calibrated, staff are bypassing safeguards or the product process is weak.

Sampling after approval remains useful, particularly for lower-risk cases that do not receive pre-approval. Samples should be risk-based and include new products, vulnerable clients, material changes, high concentrations, complex products, cross-border cases and RMs with unusual patterns.

Figure 3. Human control gates in the suitability workflow



Author framework. Material failure routes the case to clarification, escalation or rejection.

7. Validate the complete recommendation system

Validation should test the full process rather than the model alone. The system includes data sources, document retrieval, rules, calculations, prompts, model, interface, human decisions and downstream records. A technically accurate model can still produce unsuitable outcomes if client data are stale or the workflow omits total-portfolio exposure.

The validation plan should cover intended use, prohibited use, data quality, retrieval precision, extraction accuracy, calculation reproducibility, model stability, hallucination, bias, prompt

injection, access, record completeness, human factors and fallback. Tests should include ordinary cases, boundary cases and deliberate adversarial inputs.

Performance measures need decision meaning. Retrieval precision can show whether the system found approved documents. Citation accuracy can show whether statements are supported by quoted passages. False-negative rates can measure missed restrictions. Override rates can reveal workflow friction or weak controls. Outcome testing can examine whether similarly situated clients receive materially inconsistent treatment.

Acceptance thresholds are management decisions informed by obligations, harm and risk appetite. A threshold suitable for drafting a meeting summary may be inadequate for a client recommendation. Limitations should be recorded and translated into use restrictions, human checks or rejection of the use case.

Table 4. Validation and monitoring plan

Test domain	Pre-use test	Production evidence	Decision consequence
client data	completeness, freshness, reconciliation and conflict cases	missing-field rate, stale records and manual corrections	stop affected case or restrict eligible population
product evidence	approved-source retrieval, version and attribute extraction	citation accuracy, expired documents and product exceptions	remove product or refresh evidence
policy rules	positive, negative and boundary cases	rule triggers, false negatives and override patterns	correct rule, re-review cases or suspend workflow
generative output	factual support, consistency, uncertainty and prohibited claims	unsupported statement rate, edits and complaints	strengthen retrieval, prompt or human review
fairness and outcomes	comparable cases, proxies and vulnerable-client scenarios	outcome differences, declines, escalations and remediation	investigate cause and change data, process or access
security and resilience	access, prompt injection, data leakage, provider failure and fallback	alerts, incidents, downtime and recovery tests	contain, revoke access, fall back or suspend

Thresholds should be approved for the specific use case; this table defines test types rather than numerical benchmarks.

8. Govern client data as recommendation evidence

Client data are rarely a single clean record. Objectives can be expressed in meetings, questionnaires, emails and portfolio instructions. Risk tolerance may be captured through a score, while capacity for loss depends on financial commitments and liquidity. Knowledge and experience can differ across products. Legal entities, trusts and family structures can involve several decision makers.

The system should retain the source and status of each material fact. A client-stated preference is different from a verified account balance. A calculated exposure is different from an RM judgement. A missing answer is different from a negative answer. These distinctions should survive ingestion into the model.

Free-text interpretation requires confirmation. If a client says that capital should remain available for a property purchase, an AI system may classify a liquidity need. The RM should confirm

timing, amount, currency and certainty. The recorded fact should reflect what the client confirmed.

Data minimisation and access controls remain important. A model should receive only the information required for the approved task. Sensitive personal information should not be transferred to an unapproved public tool. Retention, deletion, cross-border transfer and client rights should be assessed under the applicable data framework.

9. Govern product evidence through its lifecycle

A product shelf changes continuously. Terms can be amended, subscriptions suspended, liquidity gates applied, managers replaced and risk classifications updated. Suitability systems should consume approved product evidence with clear ownership and effective dates.

The product record should cover objective, strategy, instruments, leverage, liquidity, valuation, fees, performance scenarios, loss characteristics, currency, target market, conflicts, distribution restrictions and relevant documents. Each material attribute should link to an authoritative source. Marketing text should not replace legal and approved product documents.

Product events should trigger workflow action. A suspended redemption should alter liquidity evidence. A change in manager or strategy can require a re-assessment. A document update should invalidate superseded passages. The system should identify client recommendations or holdings affected by material changes.

Complex products deserve deeper controls. Model-generated comparisons can hide path dependency, issuer credit, leverage, early termination and scenario limitations. The RM should understand the payoff and explain plausible adverse outcomes. A standard score should not substitute for that analysis.

10. Compare the recommendation with credible alternatives

Suitability is strengthened when the RM compares realistic alternatives. The set should include maintaining the current portfolio where appropriate, lower-cost or more liquid options, different implementation routes and the proposed product. The comparison should use the client's objectives and total portfolio.

AI can help identify products that meet documented filters. It can also narrow the set too aggressively through embedded assumptions or commercial priorities. Shelf coverage, ranking logic, conflicts and excluded products should therefore be visible to the RM and reviewer.

Costs should be considered cumulatively and in context. Product fees, advisory charges, transaction costs, financing costs, foreign-exchange effects and exit costs can alter the outcome. The recommendation explanation should describe the value expected from the selected option and material disadvantages relative to alternatives.

Where the firm receives placement fees, retrocessions or other benefits, the conflict should be addressed under the applicable framework. The model should not optimise commercial revenue while presenting the output as client-led. Ranking logic and overrides should be monitored for this risk.

Table 5. Hypothetical client-product comparison

Decision factor	Current diversified portfolio	Daily-liquid balanced fund	Private-market allocation	Evidence and RM question
stated objective	moderate long-term growth	broadly aligned	potentially aligned with return objective	which objective has priority and over what horizon?
liquidity	daily	daily	multi-year lock-up with limited transfer	what cash need could arise before exit?
loss capacity	hypothetical moderate capacity	within assumed range	drawdown and capital-call exposure may exceed comfort	is loss capacity verified and portfolio-wide?
concentration	current baseline	modest change	increases illiquid and manager concentration	what limit applies after the transaction?
cost	existing disclosed cost	hypothetical medium	hypothetical high with performance-related elements	which costs are certain, contingent and comparable?
operational burden	established	standard subscription	capital calls, tax and reporting complexity	can the client and firm support the lifecycle?
illustrative outcome	maintain	candidate for analysis	requires enhanced review	what evidence would change the decision?

Every value and assessment is a hypothetical management assumption used to demonstrate documentation; no product is recommended.

11. Give the RM an explanation card

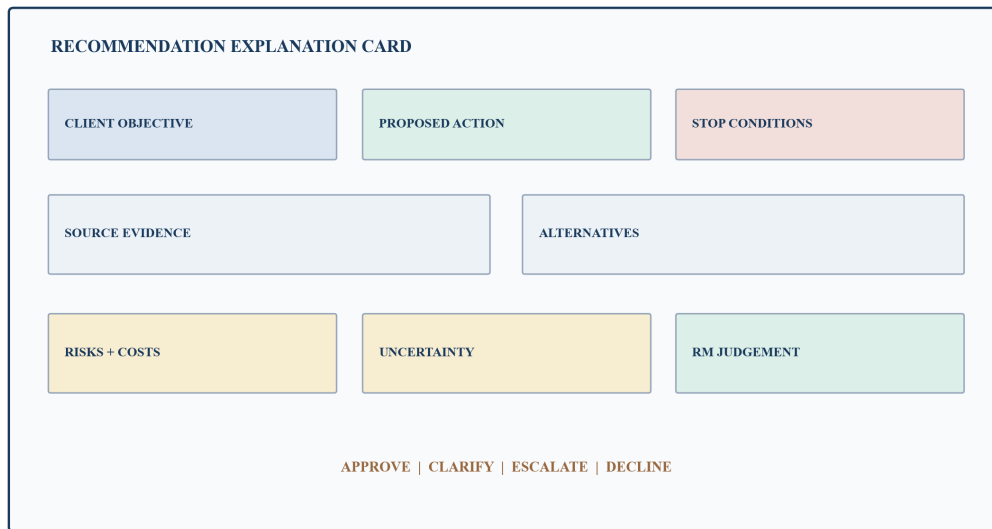
The interface should help an RM reason, not merely approve. A decision card can show the client objective, critical facts, proposed action, alternatives, supporting sources, material risks, costs, uncertainty, conflicts and required approvals. It should expose missing and stale evidence.

The card should distinguish machine-produced content. Retrieved passages can be shown with source links. Calculations can show method and inputs. Generated narrative can be editable and visibly marked within the internal workflow. The final client communication should contain only verified and approved content.

An explanation should be specific enough to test. Statements such as suitable for diversification are weak without identifying current concentration, the proposed change, time horizon, liquidity and risk. A reviewer should be able to understand how the recommendation serves the client's objective and where it could fail.

The interface should limit automation bias. The proposed product should not be visually dominant before the RM reviews evidence. Alternatives and stop conditions should receive comparable prominence. Confidence indicators should be defined and tested; a polished tone should never be treated as confidence.

Figure 4. Relationship-manager explanation card



Author framework. The internal card foregrounds evidence, uncertainty, alternatives and decision responsibility.

12. Control prompts, retrieval and model change

A generative system can change even when the user interface appears stable. The model provider can update the underlying service. The firm can change a system prompt, retrieval index, document parser, temperature, safety configuration or post-processing rule. Each material change can alter the recommendation process.

The change inventory should record component, owner, version, date, rationale, testing, approval and rollback. Regression tests should include known ordinary, boundary and failure cases. A change to a product-document parser should be tested against tables, footnotes, scanned pages and amended documents.

Prompts are control artefacts. System prompts, approved templates and prohibited instructions should be versioned and access-controlled. User prompts and model outputs should be retained where required and safe. Prompt injection and malicious document content should be included in security testing.

Third-party model updates require contractual and technical monitoring. The firm should understand notice arrangements, service logs, data use, retention, sub-processors, region, security, resilience and exit. Where a provider will not disclose model detail, the firm can still test observable performance, constrain inputs and outputs, require evidence and preserve fallback capability.

13. Design overrides as risk signals

An override can be appropriate when a rule has a documented limitation and a qualified person has better evidence. It can also be a route around a necessary control. Governance should distinguish these cases.

The override record should include the triggering condition, system output, human rationale, evidence, authority, client effect and expiry. Material overrides should receive independent

review before action. Lower-risk overrides can be sampled after action according to policy.

Patterns matter more than isolated totals. A high rate for one product may indicate poor product data. A high rate for one RM may indicate training or conduct concerns. Repeated overrides of a liquidity rule may show commercial pressure. Overrides that improve client outcomes can show a model limitation that should be corrected.

Management reporting should connect overrides to decisions. The owner should decide whether to change the rule, improve data, retrain staff, restrict the product, re-review client cases or suspend the system. Closing the ticket without testing the change leaves the risk unresolved.

14. Monitor outcomes, drift and conduct

Production monitoring should cover system performance and client outcomes. Technical measures include availability, latency, retrieval, citation accuracy, model errors and security events. Process measures include completion, stops, escalations, overrides, review time and record quality. Outcome measures include concentration, cost, complaints, cancellations, remediation and differences across client groups.

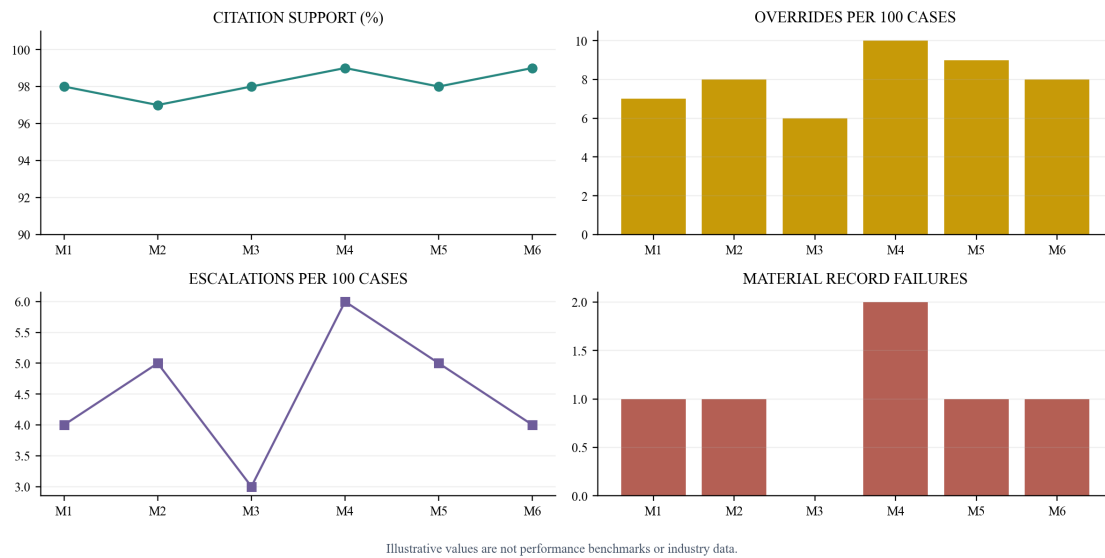
Drift can arise from market conditions, product changes, client behaviour, data sources and model updates. A stable accuracy score can hide a changing case mix. Monitoring should therefore segment by product complexity, client type, geography, channel, RM and model version where legally and operationally appropriate.

Conduct surveillance can examine whether recommendations cluster around products with higher commercial benefit, whether risk warnings are overridden, whether vulnerable clients receive different outcomes and whether generated explanations omit disadvantages. Any analysis of protected or sensitive attributes requires appropriate legal and privacy governance.

Thresholds should trigger defined actions. A critical unsupported claim can require immediate containment. A rising override rate can require enhanced sampling. A product-document failure can require removing the product from the enabled shelf. Material client impact can require case identification, communication and remediation.

Figure 5. AI suitability monitoring dashboard

HYPOTHETICAL MONTHLY CONTROL DASHBOARD



Author framework. Values are hypothetical management assumptions created to demonstrate a dashboard design.

15. Prepare for incidents and client remediation

An AI incident can involve unsupported advice, wrong product data, biased treatment, data leakage, unauthorised access, provider failure or missing records. The response plan should connect technology containment with client and regulatory decisions.

The incident lead should identify affected model versions, dates, users, clients, products and outputs. Preserved lineage makes this possible. The firm should stop or restrict the use case, secure evidence, apply fallback procedures and decide whether past recommendations require review.

Severity should reflect actual and potential client harm, regulatory breach, data sensitivity, volume and duration. Notification obligations depend on jurisdiction and facts. Legal, compliance, data protection, information security, business and communications teams need defined authority and escalation routes.

Remediation should address individual cases and systemic cause. Affected clients may require correction, explanation, portfolio action, fee adjustment or compensation according to applicable obligations and circumstances. The underlying response can require data correction, model change, rule redesign, stronger review, vendor action or use-case withdrawal.

16. Manage third-party AI as part of the regulated process

Vendor due diligence should examine the service, model, data, security, resilience, change process, audit evidence, sub-processors and contractual rights. The depth should reflect whether the provider retrieves information, ranks products, creates advice content or participates in a client decision.

The firm should understand how client and firm data are used. It should determine whether prompts or outputs train provider models, where data are processed, how long they are retained, who can access them and how deletion works. Confidentiality language should align with

technical configuration.

Operational oversight should include availability, performance, incidents, material changes and recovery. An exit plan should preserve prompts, records, product evidence, client decisions and a viable manual or alternative process. Portability matters because the recommendation record can outlive the technology contract.

Accountability should remain clear. A provider can supply a model and still lack the client context and regulated responsibility needed to approve a recommendation. The wealth firm must retain the governance and evidence needed for its own obligations.

17. Train RMs to challenge the system

Training should address the approved use case, evidence, limitations, prohibited use, client communication, overrides, incidents and fallback. It should use realistic examples where the system is correct, incomplete and confidently wrong.

RMs should be able to identify unsupported claims, stale product terms, inconsistent client facts, hidden assumptions and portfolio effects. They should understand that fluency is not evidence. A model that provides a clear explanation may still rely on the wrong source.

Competence assessment should involve observed decisions. A short knowledge test can establish understanding. Scenario work can show whether the RM challenges the tool, seeks missing facts, compares alternatives and records rationale. Repeat errors should lead to targeted coaching, restricted access or reassessment.

Managers should monitor workload. Human review loses value if approval queues are too large, evidence is difficult to navigate or commercial timing discourages challenge. Capacity is a control condition.

18. Implement the framework in 180 days

Implementation should start with an inventory of existing and planned AI uses. Shadow use through public tools should be included. Each use case should have a business owner, risk classification, data map, provider, approval status and operating restriction.

A bounded pilot should use a defined client and product population. It should test the evidence chain, control gates, RM interface, record, monitoring and fallback. Pilot results should include failed cases and staff behaviour, not only productivity.

Scaling should follow verified capability. Product coverage can expand after source governance is stable. Client coverage can expand after edge cases and fairness risks are tested. New languages and markets require their own evidence and legal analysis.

The board or delegated committee should receive an operating-state decision at the end of the programme. The pack should describe approved uses, limitations, validation results, incidents, unresolved risks, client safeguards, accountability and next review.

Table 6. 180-day implementation roadmap

Period	Primary work	Required output	Approval gate
days 1 to 30	inventory, legal perimeter, ownership and use-case classification	AI register, prohibited uses, accountable owners and evidence baseline	approve scope and pilot
days 31 to 60	client, product, policy and lineage design	source map, freshness rules, control gates and record specification	approve evidence standard
days 61 to 90	validation, security and human-factor testing	test results, limitations, fallback and remediation	permit, restrict or reject pilot use
days 91 to 120	controlled pilot and enhanced case review	outcome, override, incident and user evidence	decide product and client expansion
days 121 to 150	provider controls, monitoring and training	contracts, dashboard, competence records and response plan	accept residual risk and operating limits
days 151 to 180	independent challenge and operating-state review	assurance pack, open actions and continuing test cycle	approve production state or further restriction

Timing is an illustrative sequencing framework and should be adapted to firm size, risk and existing capability.

19. Make board reporting decision-useful

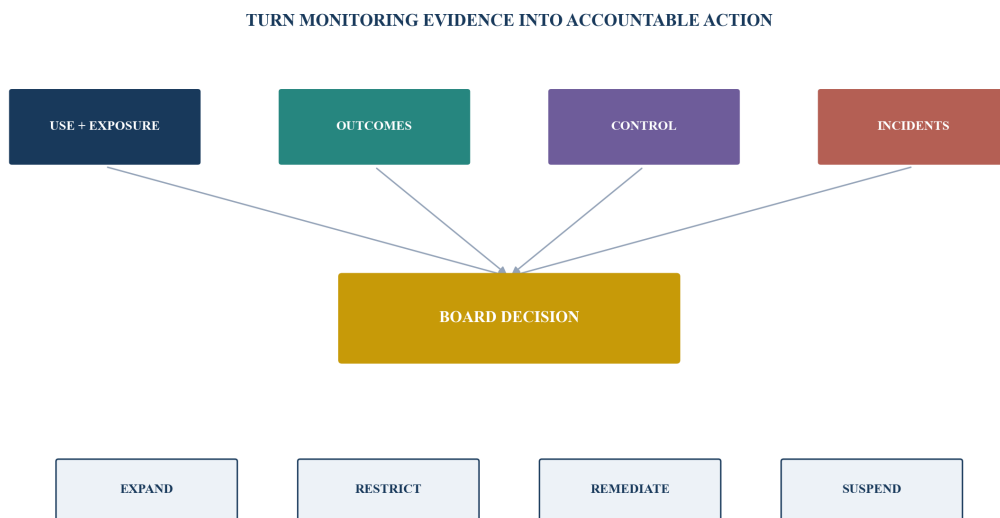
Senior management needs a view of use, exposure, performance and unresolved risk. The report should show active use cases, affected clients and products, material changes, validation status, limitations, control performance, overrides, incidents, complaints, remediation and vendor dependencies.

Metrics should be connected to thresholds and actions. A citation-support rate has value when management knows which failures affect client decisions and what response follows. An override count has value when segmented by reason, product, RM and outcome.

The board should see the operating perimeter. The pack should identify which functions are enabled, which are prohibited, which populations are excluded and which manual fallback exists. Expansion decisions should state the evidence required before approval.

Assurance should include independent challenge. Compliance monitoring, internal audit, model validation, security testing and external review can each provide different evidence. Findings should have accountable owners, deadlines and closure tests.

Figure 6. Board decision view for AI-assisted suitability



Author framework. Reporting links exposure and control evidence to explicit management decisions.

20. Make trusted advice the commercial proposition

Clients can benefit from faster preparation, more consistent evidence and clearer comparisons. Trust depends on how the firm governs the recommendation. The RM should be able to explain what information was used, what the technology contributed, which judgement the RM made and which limitations remain.

The firm should describe AI use accurately. Broad claims of objectivity or error-free automation are inappropriate. Client communication should focus on the service, safeguards, material limitations and route for questions or correction. Disclosure requirements depend on the applicable framework and facts.

A controlled suitability process can support scale. Product and client evidence become reusable, decisions become reviewable and recurring weaknesses become visible. The operating data can inform product governance, training, service design and client support.

The commercial value comes from accountable advice. A client should receive a recommendation grounded in current evidence, total-portfolio context, credible alternatives and human judgement. AI can help assemble and test that record. The firm's people remain responsible for the professional decision.

Conclusion

AI-assisted suitability should be designed as an evidence and decision system. Client facts, portfolio data, product evidence, policies, machine transformations, human judgement and client communication must remain distinct and traceable.

Practical explainability allows a qualified reviewer to reconstruct the recommendation without requiring access to proprietary model weights. The record identifies sources, versions, calculations, model contributions, alternatives, overrides, approvals and final wording. Missing and conflicting evidence remains visible.

International official materials use different legal routes. They repeatedly support governance, accountability, appropriate human involvement, data quality, testing, monitoring, provider oversight, record keeping and client protection. A wealth firm can translate those disciplines into a common operating model while completing jurisdiction-specific legal analysis.

The governing question is direct: can the relationship manager and the firm demonstrate, from current evidence, why the recommendation serves this client's documented circumstances and objectives, how the technology contributed, which alternatives were considered, which risks and costs were explained and which qualified person accepted responsibility? A complete answer turns AI from an opaque productivity tool into a controlled component of professional advice.

References

- [1] UK Financial Conduct Authority, Our approach to AI, <https://www.fca.org.uk/firms/innovation/ai-approach>
- [2] UK Financial Conduct Authority, AI in financial services: our approach, 8 June 2026, <https://www.fca.org.uk/news/blogs/ai-financial-services-approach>
- [3] UK Financial Conduct Authority, About the Consumer Duty, <https://www.fca.org.uk/firms/consumer-duty/about>
- [4] UK Financial Conduct Authority, PS22/9: A new Consumer Duty, <https://www.fca.org.uk/publications/policy-statements/ps22-9-new-consumer-duty>
- [5] UK Financial Conduct Authority, Mills Review of AI in retail financial services, 6 July 2026, <https://www.fca.org.uk/news/press-releases/fca-publishes-landmark-review-impact-ai-retail-financial-services>
- [6] European Securities and Markets Authority, Public Statement on the use of artificial intelligence in the provision of retail investment services, 30 May 2024, https://www.esma.europa.eu/sites/default/files/2024-05/ESMA35-335435667-5924_Public_Statement_on_AI_and_investment_services.pdf
- [7] Securities and Futures Commission of Hong Kong, Circular to licensed corporations on use of generative AI language models, 12 November 2024, <https://apps.sfc.hk/edistributionWeb/api/circular/list-content/circular/intermediaries/supervision/doc?lang=EN&refNo=24EC55>
- [8] Securities and Futures Commission of Hong Kong and other Hong Kong regulators, Joint Circular on the Expansion of the Generative Artificial Intelligence Sandbox, 5 March 2026, <https://apps.sfc.hk/edistributionWeb/gateway/EN/circular/intermediaries/supervision/doc?refNo=26EC11>
- [9] Australian Securities and Investments Commission, Report 798: Beware the gap, 29 October 2024, <https://asic.gov.au/regulatory-resources/find-a-document/reports/rep-798-beware-the-gap-governance-arrangements-in-the-face-of-ai-innovation/>
- [10] Australian Securities and Investments Commission, Review your artificial intelligence governance and risk management arrangements, April 2025, <https://asic.gov.au/about-asic/corporate-publications/newsletters/market-integrity-update/miu-issue-166-april-2025/>
- [11] Financial Industry Regulatory Authority, Regulatory Notice 24-09: Regulatory obligations when using generative AI and large language models, 27 June 2024, <https://www.finra.org/rules-guidance/notices/24-09>
- [12] Dubai Financial Services Authority, Artificial Intelligence Survey 2025, <https://www.dfsa.ae/news/new-df-sa-ai-survey-generative-ai-adoption-has-nearly-tripled-within-difc-last-12-months-governance-continues-develop>
- [13] Dubai Financial Services Authority, Cyber and Artificial Intelligence Risk in Financial Services, 30 June 2025, <https://www.dfsa.ae/news/new-df-sa-report-explores-regulatory-insights-cybersecurity-artificial-intelligence-and-quantum-risks>
- [14] Personal Data Protection Commission Singapore, Model Artificial Intelligence Governance Framework, second edition, <https://www.pdpc.gov.sg/-/media/files/pdpc/pdf-files/resource-for-organisation/ai/smodelai>

govframework2.pdf

- [15] Personal Data Protection Commission Singapore, Singapore's approach to AI governance, <https://www.pdpc.gov.sg/Help-and-Resources/2020/01/Model-AI-Governance-Framework>
- [16] International Organization of Securities Commissions, The use of artificial intelligence and machine learning by market intermediaries and asset managers, September 2021, <https://www.iosco.org/library/pubdocs/pdf/IOSCOPD684.pdf>

About the Author

Chennakeshav Adya is an independent researcher and Managing Partner of Matchpoint Partners. His work examines strategy, capital formation, valuation, transactions and operating execution across private and public markets.