

Sound Practices for Responsible Adoption of Artificial Intelligence (AI): Consultation report

Response to Consultation

Universal Owners Council

1. Do you agree with the benefits and risks of AI adoption by financial institutions described in this report? Are there any substantive benefits or risks not covered?

Broadly yes, subject to additional system-level, governance-chain and professional-capacity risks.

The report provides a balanced account of AI's potential to improve risk detection, operational efficiency, service quality, decision support and resilience, while addressing model performance, explainability, cyber risk, consumer harm, third-party dependence and agentic autonomy. UOC particularly supports the recognition that non-adoption may itself create risk where AI materially improves fraud detection, monitoring or operational resilience — provided this is not read by boards or supervisors as generalised adoption pressure; a deliberate, documented decision not to adopt is a legitimate outcome of the very governance the report promotes.

Seven additional risks should be made explicit:

- Accountability diffusion and “accountability laundering”. AI can create distance between the actor formally accountable for an outcome and the provider, adviser or model that materially shaped it. Vendor opacity should be linked to the broader risk that responsibility is retained but becomes practically unexercisable.
- Governance-chain conflicts. A technology provider, consultant or other intermediary may influence problem definition, model selection, configuration, validation, assurance and remediation. Where the same party performs several roles, the apparent control environment may be circular rather than independent.
- Systemic materiality and aggregation. Many low- or medium-materiality uses may share a model, dataset, cloud layer or decision logic. Their combined effect can become material even where each institution's assessment is reasonable in isolation.
- Epistemic concentration and governance monoculture. Common models and advisers can align not only trades but also the framing of risk, valuations, classifications, disclosures, covenant monitoring and escalation decisions. This may narrow institutional judgment before correlated market behaviour becomes visible.

- Erosion of professional capacity and cognitive deskilling. The report identifies automation bias and over-reliance as threats to the quality of human oversight, but not the longer-horizon version of the same problem: erosion of the sector's capacity to supervise AI at all. Analysis published by CFA Institute observes that AI output frequently appears plausible even when incorrect — which raises the cost and cognitive demand of genuine review — and highlights the need for organisational policies that guard against the deskilling of professionals. If AI absorbs the routine analytical work through which junior professionals have historically developed judgment, the apprenticeship pipeline that produces future senior validators, risk-takers and overseers weakens across the sector simultaneously. The sound practices assume a persistent supply of humans capable of meaningful oversight; that assumption should itself be risk-managed.

- Loss of recovery capacity under prolonged AI outages. The report's business-continuity provisions assume that AI-enabled services can continue manually in the event of disruption. Where efficiency gains have been captured by reducing the teams that previously performed the work, manual fallback becomes a paper control: the process documentation survives, but the people who can execute it at volume, under stress, do not. The risk is correlated: a prolonged outage at a shared foundation model or cloud provider affects many institutions simultaneously, precisely when each needs surge human capacity and none can hire or retrain within the disruption window. Recoverability erodes gradually and invisibly with each year of AI dependence, with nothing observable until the outage occurs. Continuity plans that assert manual continuation without testing it against actual staffing, skills and throughput are untested liabilities.

- Contestability and redress failure. Where AI affects credit, pricing, insurance, investment, valuation, employment, market access or legal rights, affected stakeholders need a practical ability to challenge the outcome, obtain meaningful review and secure correction or redress. Explainability without contestability may not protect institutional legitimacy.

2. Are the sound practices sufficiently comprehensive and clear to enable financial institutions' responsible AI adoption?

They are comprehensive as a foundation, but several control outcomes should be stated more clearly.

The framework identifies the right domains. The final report should sharpen the expected governance architecture for high-materiality, high-autonomy or systemically connected uses. In particular:

- Separate governance roles. Distinguish the business owner, system operator, independent validator, approval authority, second-line risk function and internal audit or assurance role. Approval stages should separate operators, validators and approvers; departures should be documented and justified.

- Make accountability named and non-delegable. Each material AI use should have an accountable executive and identified business owner. Outsourcing development or operation should not outsource accountability.

- Anchor oversight in individual professional accountability. The sound practices allocate accountability to organisational roles; they say little about the individual professionals who will exercise the oversight the practices demand. A mature accountability infrastructure for individuals already exists: the CFA Institute Code of Ethics and Standards of Professional Conduct, binding on CFA Institute members and candidates globally, requires individual competence, diligence and reasonable basis, avoidance of misrepresentation, and conduct preserving the integrity of capital markets, and CFA Institute guidance applies these obligations directly to AI-assisted work — including expectations that models be periodically tested so that automated decision-making does not produce manipulative or market-distorting outcomes. Individual professional accountability has a property organisational controls lack: it survives even where model explainability does not. A professional who signs work product remains answerable for it whether or not the model that assisted them can be interrogated. The final report should acknowledge professional-conduct frameworks as complementary infrastructure, and Sound Practice 2 should encourage institutions to align internal accountability with the professional obligations their staff already carry.
- Treat competence maintenance as a control. Sound Practice 4 addresses AI literacy and role-specific training; Sound Practice 10 addresses incentives for thoughtful review. Neither addresses preservation of the underlying analytical skill that makes review meaningful. CFA Institute’s ongoing work on the structural implications of AI for the profession frames the durable model as the pairing of artificial with human intelligence — augmentation, not substitution — precisely because professional judgment is required to test assumptions and validate what models produce. Training programmes should include periodic unassisted exercise of core analytical skills for staff in oversight roles, and institutions should monitor for skill atrophy as deliberately as they monitor for model drift. A control function that can no longer perform the analysis it is checking is a control in name only.
- Require evidence, not only process. Governance should generate an auditable evidence package covering purpose, owner, data lineage, provider dependencies, materiality assessment, validation record, approval basis, performance thresholds, incidents, material changes, override decisions and exit — capable of supporting external assurance and appropriate investor-facing disclosure, not only internal risk management.
- Protect independent escalation. Employees, contractors and relevant third parties should be able to escalate unsafe outputs, undisclosed changes, control weaknesses and near misses without interference or retaliation. Boards or delegated committees should receive material themes and unresolved disagreements.
- Distinguish board and management responsibilities. The board should approve appetite, prohibited uses and governance boundaries; management should implement; independent risk should challenge and validate; internal audit should assess design and effectiveness. Combining “board and senior management” too frequently risks blurring distinct decision functions.
- Introduce an ecosystem layer. A companion principle should address shared providers, common models, pooled assurance, sector exercises and supervisory aggregation. An institution-by-institution framework cannot fully manage financial-stability externalities.

Finally, the report would be clearer if it stated explicitly — rather than leaving it to be inferred from “all types of financial institutions” — that asset owners, including pension funds, sovereign wealth funds and endowments, are within scope.

3. Do the sound practices strike an appropriate balance between managing risks relating to all forms of AI, and addressing some of the risks relating to emerging and new complex forms of AI, such as GenAI and agentic AI?

Broadly yes. The final report should remain technology-neutral but become more capability- and decision-rights-sensitive.

The relevant trigger is not the label attached to a technology. It is what the system can do, the authority it can exercise, the systems it can reach, the speed and scale at which it can act, and the reversibility of its actions. Controls should therefore be calibrated to autonomy, tool use, external connectivity, memory, self-modification, permissions, transaction execution and multi-agent interaction.

For agentic AI, the final report should state a stronger decision architecture:

- No self-approval. An agent should not originate, validate and authorise the same material action.
- Separation of material actions. Recommendation, validation, authorisation and execution should be appropriately separated for material financial, legal, customer or market actions.
- Controlled permissions. Agents should not expand their own permissions, change control thresholds or disable monitoring without independent authorisation.
- Identity and bounded authority. Material actions should have explicit transaction limits, identity, logging, revocation, suspension and recovery arrangements.
- Substantive, not ceremonial, oversight. Where meaningful human review cannot occur at machine speed, the answer should be constrained autonomy, pre-authorised boundaries, ex-post surveillance and automatic suspension triggers rather than nominal real-time review.

Beyond any single institution, the interaction between autonomous agents deployed by different institutions in the same markets is a channel through which herding, liquidity withdrawal or manipulative dynamics could propagate at machine speed with no single institution at fault. CFA Institute guidance made an early version of this point in the pre-agentic context, noting the importance of the ability to override algorithms in systemic events where machine-learning systems used by different market participants may act in concert to amplify stress. Cross-institution agent interaction should be flagged as a topic for future FSB financial-stability work, since it sits definitionally beyond any one institution’s risk management.

4. Are the sound practices sufficiently flexible to accommodate and address newer types of AI and responsible adoption over time?

Yes in structure, provided flexibility is anchored to stable control outcomes and mandatory reassessment triggers.

The principles-based structure is durable. Flexibility should not, however, permit controls to drift as providers substitute models, activate new AI features or expand agent permissions. Reassessment should follow a material change, including changes to the underlying model, training or retrieval data, guardrails, prompts, APIs, tools, permissions, hosting location, subcontractors, performance profile or intended use.

A durable architecture would combine:

- Stable governance outcomes. Accountability, independence, traceability, security, resilience, contestability and redress.
- Capability-based control triggers. Autonomy, scale, connectivity, speed, explainability, reversibility and access to external tools or systems.
- Event-driven reassessment. Material changes, incidents, near misses, drift, provider changes, prolonged provider outages, new legal obligations and changes in aggregate sector exposure.

Scenario testing should evolve with dependence. Sector exercises — building on the CMORG-style scenario libraries the report cites — should include the prolonged unavailability of a widely used model, cloud or data provider, testing not only technical failover but the residual human capacity of many institutions to operate critical processes manually at the same time.

The FSB should consider a periodic implementation note or living annex — including a maintained glossary — updated through supervisory and industry evidence, rather than reopening the core principles whenever technology changes.

5. Do the case studies in this report sufficiently highlight how different types of financial institutions can benefit from responsible AI adoption? Are there additional case studies for inclusion in the report? If so, please provide such case studies, particularly for nonbanks.

The case studies are useful but remain weighted toward large banks and insurers. Additional nonbank and ecosystem cases are needed.

The Bank of England/FCA survey indicates broad adoption across financial services, while IOSCO and IAIS now provide sector-specific supervisory work. The final report should use those differences to test how the sound practices operate across asset management, private markets, long-horizon asset ownership, financial market infrastructure and outsourced administration. No pension fund, sovereign wealth fund or endowment appears in the current draft, despite the report's stated application to all types of financial institutions.

UOC recommends at least four additional case studies, and would welcome the opportunity to convene asset owners and relevant experts to help develop them:

1. Asset management: common-model liquidity and execution risk. Several asset managers use the same third-party model for liquidity scoring, portfolio rebalancing or execution. A model update changes risk estimates across firms and produces synchronised trading. The case should show dependency mapping, independent benchmarks, material-change notification, portfolio limits, manual fallback and a cross-market scenario exercise.

2. Private credit and private markets: embedded AI across the advisory chain. A manager uses AI embedded in underwriting, valuation, covenant monitoring, fund administration and legal-document workflows. The case should address conflicts where one intermediary configures and validates the system, the rights of limited partners and boards to obtain evidence, independent review, version control, escalation and the effect of common service providers across multiple funds.

3. Pension fund or sovereign investor: governance purchasing power. A long-horizon owner procures AI-enabled investment and advisory services using standard disclosure, audit, incident-notification, portability and model-change clauses. Multiple owners use pooled assurance and common vendor questionnaires. The case would demonstrate how repeat purchasers can improve market standards rather than act only as passive recipients of provider terms.

4. Financial market infrastructure or fund administrator: operational continuity under prolonged AI failure. An AI-enabled process affects collateral, settlement, reconciliation, transfer agency or records. The case should test severe but plausible prolonged failure, fourth-party dependency, reconciliation controls, manual continuation against actual staffing and throughput, data portability and coordinated incident response.

6. Do the case studies in this report provide actionable insights for financial institutions in their responsible AI adoption?

Partly. They illustrate practices, but a common case-study template — and methodological caveats on reported metrics — would make the insights more operational.

Each case study should state, in a consistent format:

- Decision and stakeholders. The decision or service being supported and the affected stakeholders.
- Accountability map. The accountable owner and the roles of operator, validator, approver and assurer.
- Capability profile. The AI capability, degree of autonomy and external systems or data it can access.
- Materiality. The assessment at use-case, institution, portfolio and system levels.
- Dependency map. Critical third-, fourth- and nth-party dependencies.
- Preventive controls. Pre-deployment controls, performance thresholds and conditions for suspension.

- Evidence and escalation. Records retained, material-change triggers and the escalation pathway.
- Resilience and redress. Failure scenario, fallback, exit and redress arrangements.
- Learning. Lessons transmitted to the board, supervisors and peer institutions.

This format would help institutions translate narrative examples into decision-rights matrices and control designs, and improve comparability across sectors and jurisdictions without turning the sound practices into prescriptive rules.

Reported metrics should carry a standing caveat. Several case studies quote striking outcome figures — detection rates above ninety-five per cent, fraud-loss reductions above twenty per cent, turnaround times cut from days to seconds. These are useful illustrations, but unaudited firm-reported metrics reproduced in an FSB document acquire a weight their provenance may not support, and risk anchoring board and supervisory expectations. A short statement that figures are self-reported and context-dependent, together with encouragement to disclose evaluation methodology alongside results, would preserve the illustrative value while protecting the report’s authority.

7. Are the definitions in the glossary clear and aligned with industry sound practices, including recent developments in AI?

The existing definitions are generally clear. Several governance and ecosystem terms should be added or refined, and terms used in the body text should be promoted to the glossary.

Terms used in the report but absent from the glossary. “Shadow AI” appears in the body text (Sections 1.3 and 3.8) and is a concept boards will need to operationalise, yet it is not defined in the glossary. “Reward hacking” and “automation bias” are defined only in footnotes despite carrying substantive weight in the agentic-risk and human-oversight discussions. The taxonomy of oversight modes in Sound Practice 10 (human-in-the-loop, human-on-the-loop, AI-in-the-loop, human-in-command) is likewise glossary-worthy, since these terms are used inconsistently across the industry. All should be promoted to the glossary.

Recommended additions:

- AI agent. An AI system capable of pursuing goals through a sequence of actions, including the use of tools, data, APIs, systems or other agents, within defined or adaptive permissions.
- Human oversight. A control function performed by one or more competent and authorised persons who have sufficient information, independence, time and practical ability to challenge, intervene, suspend, reverse or escalate an AI-supported action.
- Material change. A change to a model, system, data, guardrail, prompt, tool, permission, provider, subcontractor, hosting arrangement or use case that could alter performance, risk, explainability, resilience or legal obligations.

- **Embedded AI.** AI functionality incorporated into a wider product or service, including professional, operational, investment, legal, audit, data or technology services, whether or not marketed as an AI service.
- **Systemic materiality.** The potential for an AI use or dependency to affect financial stability through scale, interconnectedness, common adoption, substitutability constraints, correlated outcomes or concentration, including where individual uses are not material in isolation.
- **Independent validation.** Effective challenge by a function sufficiently separate from the design, operation and commercial sponsorship of the AI use, with the information, competence, organisational standing and authority required to assess fitness for purpose and control effectiveness.
- **Contestability.** The ability of an affected person or institution to obtain meaningful review of an AI-supported outcome, challenge relevant evidence or assumptions, and secure correction or redress where appropriate.
- **AI incident and near miss.** An event, control failure or abnormal behaviour that caused, or could reasonably have caused, material harm, breach, disruption or unauthorised action.
- **Supervisory observability.** The capacity of relevant authorities to identify and aggregate material AI uses, dependencies, providers, incidents and concentration across institutions sufficiently to assess system-wide risk.

The definition of “supply chain” should also be reconsidered. Limiting it to services within a direct third-party relationship may be too narrow for AI ecosystems in which dependencies are indirect, embedded, open-source, vertically integrated or introduced through subcontractors. The final definition should capture material indirect dependencies while preserving proportionality. Finally, the glossary or introduction should state expressly that “financial institution” includes asset owners such as pension funds, sovereign wealth funds and endowments.

8. Are there any comments you would like to make on this report? Please fill in.

Universal Owners Council Limited (UOC) has prepared a detailed written response addressing governance-chain integrity, multi-level materiality, effective human oversight, professional capacity, critical recovery capacity, embedded and indirect third-party AI, investor stewardship and system-wide resilience.

UOC’s complete 23-page response, entitled Sound Practices for Responsible Adoption of Artificial Intelligence — Response to the Financial Stability Board Consultation Report of 10 June 2026, is being sent separately to fsb@fsb.org as supplementary material in accordance with the consultation instructions. Please read the emailed PDF together with this online-form submission.

UOC consents to publication of its response.

Contact: uoc@uocouncil.org

CONSULTATION RESPONSE

Sound Practices for Responsible Adoption of Artificial Intelligence

Response to the Financial Stability Board consultation report of 10 June 2026

Core proposition. Responsible AI in finance requires more than model controls and a nominal “human in the loop”. It requires an institution in command: a traceable, independent and contestable decision architecture in which practical control remains aligned with formal accountability — and a profession that retains the competence, individual accountability and standing to exercise that control.

Submitted by: Universal Owners Council Limited (UOC)

Respondent profile: UK company limited by guarantee; non-profit, non-partisan and non-commercial public-interest fiduciary platform

Contact e-mail: uoc@uocouncil.org

Submission date: 17 July 2026

Publication status: This response may be published

Purpose and institutional basis

UOC supports the FSB's effort to establish a coherent, technology-neutral foundation for responsible AI adoption. This submission focuses on the additional governance, accountability and market-structure safeguards required where AI capability is distributed across financial institutions, technology providers and professional intermediaries, and where individually rational deployment choices may create common system-wide dependencies. [1, 2, 3]

The submission applies UOC's public-interest fiduciary mandate: long-horizon owners are exposed not only to portfolio-level outcomes, but also to the integrity, resilience and legitimacy of the institutional systems on which diversified returns depend. The UOC Charter therefore emphasises systemic stewardship, independence from conflicted advisory influence, transparency, accountability, collective responsibility and protected fiduciary inquiry. [21]

The external authorities cited below are used as interoperability benchmarks. They differ in legal status, institutional scope and jurisdictional application; UOC does not suggest that each instrument applies directly to every financial institution. Their convergence is nevertheless relevant to the FSB's principles-based framework. [5, 7, 8, 16, 17, 19, 20, 22]

Executive summary

UOC broadly supports the 12 proposed sound practices. The framework is well structured, proportionate and appropriately attentive to governance, lifecycle controls, cyber resilience and third-party concentration. Its principal limitation is that it remains predominantly institution-centric, while several of the most consequential AI risks are ecosystem risks. [1, 2, 3]

AI risk in finance is not only model risk or technology risk. It is also governance-chain and market-structure risk. The same cloud platform, foundation model, data source, systems integrator, asset manager, consultant, auditor, administrator or legal adviser may shape decisions across many institutions. Formal accountability may remain with a financial institution even as the information, capability and practical control needed to discharge that accountability migrate outside it. This produces an **accountability paradox**: responsibility is retained, but control is dispersed. [3, 4, 5, 6, 7, 10, 14]

The final report would be materially strengthened by eight refinements:

- 1. Protect the integrity of the governance chain.** For material AI uses, distinguish business ownership, operation, independent validation, approval, second-line challenge and assurance. A party that selects or configures a material AI system should not also provide its sole validation or assurance without transparent safeguards and independent challenge. [7, 9, 11, 12, 17, 19]
- 2. Apply multi-level materiality.** Assess materiality at use-case, institution, group, portfolio and financial-system levels. A use that is individually modest may become systemically important through common adoption, correlated outputs, shared data or infrastructure, or cumulative exposure. [2, 3, 12, 14]
- 3. Define effective human oversight as a control capability.** Oversight is effective only where responsible persons have sufficient competence, information, time, independence, authority and operational ability to intervene, suspend, reverse or escalate. A nominal reviewer is not a control. [8, 16, 17, 20]
- 4. Sustain the professional capacity on which oversight depends.** Effective oversight presupposes professionals who remain competent, independent-minded and individually accountable. Research and guidance from CFA Institute emphasise that AI should complement human intelligence rather than substitute for professional judgment, warn that plausible-but-incorrect AI output raises the cost and cognitive demand of genuine review, and highlight the need for organisational policies that guard against cognitive deskilling. Individual professional-conduct frameworks provide accountability infrastructure that survives even where model explainability does not, and should be recognised as complementary to the sound practices. [22, 23, 24, 25, 26]
- 5. Maintain critical redundancies and recovery capacity as a condition of reliance.** Efficiency gains from AI should not be captured by eroding the human capacity, substitutable arrangements, independent data access and rehearsed manual procedures on which recovery from prolonged AI outages depends. Manual fallback that exists only in continuity documentation is not a control. Because competitive pressure will otherwise underprovide costly redundancy — the same market failure that justifies capital and liquidity buffers — defined minimum redundancies for critical operations should be

treated as a condition of reliance on AI, not an object of cost-benefit optimisation. [1, 4, 14, 15]

- 6. Extend third-party governance to embedded and indirect AI.** Sound Practice 12 should cover AI embedded in professional, investment, legal, audit, administrative, data and operational services; material model substitutions; fourth- and nth-party dependencies; conflicts of role; portability; and exit. [4, 5, 6, 7, 10, 15]
- 7. Use collective purchasing power, supervisory observability and investor stewardship.** Common contractual standards, pooled assurance, interoperable evidence, shared incident learning and sector-wide dependency mapping can convert fragmented buyer dependence into collective market discipline. Disclosure sufficient for investors and asset owners to assess an institution's AI governance harnesses stewardship as a further discipline mechanism. Institution-level controls alone cannot resolve system-wide concentration. [2, 4, 5, 14, 15, 21]
- 8. Protect escalation and organisational learning.** Incidents, near misses, undisclosed changes and unresolved validation disagreements should be capable of reaching independent control functions and boards through protected channels, with lessons fed back into risk appetite, procurement and supervisory monitoring. Professional codes of conduct reinforce this: they give individuals independent standing to escalate against commercial pressure. [9, 11, 17, 21, 22]

Priority recommendations

Area	Recommended refinement	Policy objective
Governance and accountability	Document distinct roles for owner, operator, validator, approver and assurer; preserve independent escalation.	Prevent circular assurance, role conflict and accountability diffusion.
Materiality	Add cumulative and system-level materiality alongside use-case materiality.	Capture common dependencies, correlated outputs and aggregate lower-risk uses.
Agentic AI	Separate recommendation, validation, authorisation and execution; prohibit self-approval and uncontrolled permission expansion.	Constrain autonomous action and preserve decision rights.
Human capacity and professional accountability	Preserve analytical competence in oversight roles; monitor skill atrophy; align internal accountability with professional-conduct obligations.	Ensure the supply of humans capable of meaningful oversight does not erode as AI absorbs analytical work.
Recovery capacity and redundancy	Maintain and test minimum human, provider, data and procedural redundancies for critical operations; include prolonged-outage scenarios in sector exercises.	Ensure manual fallback is a real capability, not a documentation artefact.
Third parties	Cover embedded AI, indirect dependencies, material-change notification, audit, incident, portability and exit rights.	Align practical control with retained institutional accountability.
Transparency to investors	Extend Sound Practice 8 transparency explicitly to investors and asset owners.	Enable stewardship as a complement to supervision.
Market-wide resilience	Encourage common procurement standards, pooled assurance, sector exercises and supervisory exposure mapping.	Address concentration and governance monoculture at source.
Escalation and learning	Protect channels for incidents, near misses and control concerns, with feedback to boards and risk functions.	Turn audit trails into accountability and organisational learning.

Policy basis and UOC's original contribution

The recommendations below do not depend on a single framework. They reflect convergence across financial-stability analysis, prudential model-risk management, securities and insurance supervision, third-party resilience, cross-sector AI governance, emerging legislation and professional standards. UOC's contribution is to connect those established principles to the governance and market structure of a financial system increasingly dependent on common AI providers and professional intermediaries. [3, 5, 7, 8, 11, 12, 16, 17, 19, 20, 22]

Policy issue	Established authority	UOC extension
Governance and validation	PRA and US interagency guidance support clear accountability, independent validation and effective challenge; BIS and NIST support lifecycle governance and three lines of defence.	Apply role separation across the full AI governance chain, including advisers and providers, and test whether assurance is genuinely independent.
Materiality	FSB and Bank of England analysis address concentration and correlated behaviour; US guidance recognises aggregate model risk and common dependencies.	Require multi-level materiality across use cases, portfolios, groups and the financial system, including cumulative exposure.
Human oversight	OECD, NIST, IAIS and the EU AI Act support human-centred accountability, oversight, monitoring and redress.	Define oversight by practical intervention capability rather than the mere presence of a human reviewer.
Professional integrity	CFA Institute Code and Standards bind individual professionals to competence, diligence and market integrity; CFA research applies these duties to AI-assisted work and warns of cognitive deskilling.	Treat professional capacity as a control in its own right: preserve analytical skill, individual accountability and the apprenticeship pipeline that produces future overseers.
Third-party AI	FSB, BCBS and IOSCO frameworks broaden attention from outsourcing to third-party arrangements, supply chains and concentration.	Capture embedded AI within professional services and conflicts where one intermediary performs multiple governance roles.
System-wide resilience	FSB monitoring, Bank of England macroprudential work and the UK CTP regime support aggregated visibility and system-level testing.	Create supervisory observability across institutions and use pooled assurance, procurement and investor stewardship to influence provider conduct.
Escalation and legitimacy	Risk frameworks require documentation, challenge and governance reporting; IAIS and EU rules emphasise redress.	Protect escalation by staff, contractors and third parties and connect near misses to board learning and institutional legitimacy.

UOC's central test: formal accountability is credible only where the accountable institution retains, or can reliably obtain, the information, authority and practical capability required to exercise control — and where the individuals exercising that control remain competent and personally answerable for it.

Overall assessment

The Consultation Report correctly identifies four vulnerabilities with particular financial-stability relevance: third-party dependencies and concentration, market correlations, cyber risk, and model risk, data quality and governance. It also recognises that providers of cloud infrastructure, hardware and foundation models are concentrated and sometimes vertically integrated, and that common models, data or infrastructure may produce homogeneous behaviour and correlated outcomes. [1, 3, 6, 14]

The report's 12 sound practices are therefore a strong foundation. However, the proposed controls are generally framed as actions taken by each financial institution. That framing is necessary but not sufficient. Common dependencies are not fully visible from the perspective of any one buyer; a single institution may lack leverage to obtain meaningful transparency, auditability or portability; and a provider's systemic importance may emerge only when exposures are aggregated across the sector. [2, 4, 5, 14, 15]

The final report should distinguish three layers of responsibility:

- **Institutional responsibility.** Governance, due diligence, controls, monitoring, challenge, contingency and redress within each financial institution. [1, 7, 8, 11, 17]
- **Collective market responsibility.** Common assurance standards, pooled testing, interoperable evidence, incident learning and purchasing coordination among institutions with shared dependencies — including the stewardship of asset owners as repeat purchasers and as investors in financial institutions. [4, 5, 15, 21]
- **Public-authority responsibility.** Cross-sector exposure mapping, systemic scenario analysis, supervisory coordination and, where legal mandates permit, direct oversight or equivalent assurance for critical AI-enabled services. [2, 3, 14, 15]

Responses to consultation questions

Question 1. Benefits and risks

Do you agree with the benefits and risks of AI adoption by financial institutions described in the report? Are there any substantive benefits or risks not covered?

Response: Broadly yes, subject to additional system-level, governance-chain and professional-capacity risks.

The report provides a balanced account of AI’s potential to improve risk detection, operational efficiency, service quality, decision support and resilience, while addressing model performance, explainability, cyber risk, consumer harm, third-party dependence and agentic autonomy. UOC particularly supports the recognition that non-adoption may itself create risk where AI materially improves fraud detection, monitoring or operational resilience — provided this is not read by boards or supervisors as generalised adoption pressure; a deliberate, documented decision not to adopt is a legitimate outcome of the very governance the report promotes. [1, 3, 6, 13, 14]

Seven additional risks should be made explicit:

- **Accountability diffusion and “accountability laundering”.** AI can create distance between the actor formally accountable for an outcome and the provider, adviser or model that materially shaped it. Vendor opacity should be linked to the broader risk that responsibility is retained but becomes practically unexercisable. [5, 7, 8, 10]
- **Governance-chain conflicts.** A technology provider, consultant or other intermediary may influence problem definition, model selection, configuration, validation, assurance and remediation. Where the same party performs several roles, the apparent control environment may be circular rather than independent. [7, 11, 12]
- **Systemic materiality and aggregation.** Many low- or medium-materiality uses may share a model, dataset, cloud layer or decision logic. Their combined effect can become material even where each institution’s assessment is reasonable in isolation. [2, 3, 12, 14]
- **Epistemic concentration and governance monoculture.** Common models and advisers can align not only trades but also the framing of risk, valuations, classifications, disclosures, covenant monitoring and escalation decisions. This may narrow institutional judgment before correlated market behaviour becomes visible. [3, 7, 14]
- **Erosion of professional capacity and cognitive deskilling.** The report identifies automation bias and over-reliance as threats to the quality of human oversight, but not the longer-horizon version of the same problem: erosion of the sector’s capacity to supervise AI at all. Analysis published by CFA Institute observes that AI output frequently appears plausible even when incorrect — which raises the cost and cognitive demand of genuine review — and highlights the need for organisational policies that guard against the deskilling of professionals. If AI absorbs the routine analytical work through which junior professionals have historically developed judgment, the apprenticeship pipeline that produces future senior validators, risk-takers and overseers weakens across the sector simultaneously. The sound practices assume a persistent supply of humans capable of meaningful oversight; that assumption should itself be risk-managed. [22, 25, 26]

- **Loss of recovery capacity under prolonged AI outages.** The report's business-continuity provisions assume that AI-enabled services can continue manually in the event of disruption. Where efficiency gains have been captured by reducing the teams that previously performed the work, manual fallback becomes a paper control: the process documentation survives, but the people who can execute it at volume, under stress, do not. The risk is correlated: a prolonged outage at a shared foundation model or cloud provider affects many institutions simultaneously, precisely when each needs surge human capacity and none can hire or retrain within the disruption window. Recoverability erodes gradually and invisibly with each year of AI dependence, with nothing observable until the outage occurs. Continuity plans that assert manual continuation without testing it against actual staffing, skills and throughput are untested liabilities. [1, 4, 14, 15]
- **Contestability and redress failure.** Where AI affects credit, pricing, insurance, investment, valuation, employment, market access or legal rights, affected stakeholders need a practical ability to challenge the outcome, obtain meaningful review and secure correction or redress. Explainability without contestability may not protect institutional legitimacy. [7, 8, 16, 20]

Question 2. Comprehensiveness and clarity

Are the sound practices sufficiently comprehensive and clear to enable financial institutions' responsible AI adoption?

Response: They are comprehensive as a foundation, but several control outcomes should be stated more clearly.

The framework identifies the right domains. The final report should sharpen the expected governance architecture for high-materiality, high-autonomy or systemically connected uses. In particular:

- **Separate governance roles.** Distinguish the business owner, system operator, independent validator, approval authority, second-line risk function and internal audit or assurance role. Approval stages should separate operators, validators and approvers; departures should be documented and justified. [9, 11, 12, 17]
- **Make accountability named and non-delegable.** Each material AI use should have an accountable executive and identified business owner. Outsourcing development or operation should not outsource accountability. [1, 5, 7, 8, 19]
- **Anchor oversight in individual professional accountability.** The sound practices allocate accountability to organisational roles; they say little about the individual professionals who will exercise the oversight the practices demand. A mature accountability infrastructure for individuals already exists: the CFA Institute Code of Ethics and Standards of Professional Conduct, binding on CFA Institute members and candidates globally, requires individual competence, diligence and reasonable basis, avoidance of misrepresentation, and conduct preserving the integrity of capital markets, and CFA Institute guidance applies these obligations directly to AI-assisted work — including expectations that models be periodically tested so that automated decision-making does not produce manipulative or market-distorting outcomes. Individual professional accountability has a property organisational controls lack: it survives even where model explainability does not. A professional who signs work product remains answerable for it whether or not the model that assisted them can be interrogated. The final report should acknowledge professional-conduct frameworks as complementary infrastructure, and Sound Practice 2 should encourage institutions to align internal accountability with the professional obligations their staff already carry. [22, 23, 24]
- **Treat competence maintenance as a control.** Sound Practice 4 addresses AI literacy and role-specific training; Sound Practice 10 addresses incentives for thoughtful review. Neither addresses preservation of the underlying analytical skill that makes review meaningful. CFA Institute's ongoing work on the structural implications of AI for the profession frames the durable model as the pairing of artificial with human intelligence — augmentation, not substitution — precisely because professional judgment is required to test assumptions and validate what models produce. Training programmes should include periodic unassisted exercise of core analytical skills for staff in oversight roles, and institutions should monitor for skill atrophy as deliberately as they monitor for model drift. A control function that can no longer perform the analysis it is checking is a control in name only. [24, 25, 26]

- **Require evidence, not only process.** Governance should generate an auditable evidence package covering purpose, owner, data lineage, provider dependencies, materiality assessment, validation record, approval basis, performance thresholds, incidents, material changes, override decisions and exit — capable of supporting external assurance and appropriate investor-facing disclosure, not only internal risk management. [11, 12, 17, 19, 20]
- **Protect independent escalation.** Employees, contractors and relevant third parties should be able to escalate unsafe outputs, undisclosed changes, control weaknesses and near misses without interference or retaliation. Boards or delegated committees should receive material themes and unresolved disagreements. [9, 11, 17, 21]
- **Distinguish board and management responsibilities.** The board should approve appetite, prohibited uses and governance boundaries; management should implement; independent risk should challenge and validate; internal audit should assess design and effectiveness. Combining “board and senior management” too frequently risks blurring distinct decision functions. [1, 8, 9, 11]
- **Introduce an ecosystem layer.** A companion principle should address shared providers, common models, pooled assurance, sector exercises and supervisory aggregation. An institution-by-institution framework cannot fully manage financial-stability externalities. [2, 4, 5, 14, 15]

Finally, the report would be clearer if it stated explicitly — rather than leaving it to be inferred from “all types of financial institutions” — that asset owners, including pension funds, sovereign wealth funds and endowments, are within scope.

Question 3. Balance across forms of AI

Do the sound practices strike an appropriate balance between managing risks relating to all forms of AI and addressing emerging and complex forms such as generative and agentic AI?

Response: Broadly yes. The final report should remain technology-neutral but become more capability- and decision-rights-sensitive.

The relevant trigger is not the label attached to a technology. It is what the system can do, the authority it can exercise, the systems it can reach, the speed and scale at which it can act, and the reversibility of its actions. Controls should therefore be calibrated to autonomy, tool use, external connectivity, memory, self-modification, permissions, transaction execution and multi-agent interaction. [1, 17, 18, 20]

For agentic AI, the final report should state a stronger decision architecture:

- **No self-approval.** An agent should not originate, validate and authorise the same material action.
- **Separation of material actions.** Recommendation, validation, authorisation and execution should be appropriately separated for material financial, legal, customer or market actions. [11, 12]
- **Controlled permissions.** Agents should not expand their own permissions, change control thresholds or disable monitoring without independent authorisation. [17, 18, 20]
- **Identity and bounded authority.** Material actions should have explicit transaction limits, identity, logging, revocation, suspension and recovery arrangements. [1, 18, 20]
- **Substantive, not ceremonial, oversight.** Where meaningful human review cannot occur at machine speed, the answer should be constrained autonomy, pre-authorised boundaries, ex-post surveillance and automatic suspension triggers rather than nominal real-time review. [1, 8, 16, 20]

Beyond any single institution, the interaction between autonomous agents deployed by different institutions in the same markets is a channel through which herding, liquidity withdrawal or manipulative dynamics could propagate at machine speed with no single institution at fault. CFA Institute guidance made an early version of this point in the pre-agentic context, noting the importance of the ability to override algorithms in systemic events where machine-learning systems used by different market participants may act in concert to amplify stress. Cross-institution agent interaction should be flagged as a topic for future FSB financial-stability work, since it sits definitionally beyond any one institution's risk management. [3, 14, 23]

Question 4. Flexibility over time

Are the sound practices sufficiently flexible to accommodate and address newer types of AI and responsible adoption over time?

Response: Yes in structure, provided flexibility is anchored to stable control outcomes and mandatory reassessment triggers.

The principles-based structure is durable. Flexibility should not, however, permit controls to drift as providers substitute models, activate new AI features or expand agent permissions. Reassessment should follow a material change, including changes to the underlying model, training or retrieval data, guardrails, prompts, APIs, tools, permissions, hosting location, subcontractors, performance profile or intended use. [5, 17, 18, 19, 20]

A durable architecture would combine:

- **Stable governance outcomes.** Accountability, independence, traceability, security, resilience, contestability and redress. [16, 17, 19, 20]
- **Capability-based control triggers.** Autonomy, scale, connectivity, speed, explainability, reversibility and access to external tools or systems. [1, 18, 20]
- **Event-driven reassessment.** Material changes, incidents, near misses, drift, provider changes, prolonged provider outages, new legal obligations and changes in aggregate sector exposure. [2, 5, 17, 19]

Scenario testing should evolve with dependence. Sector exercises — building on the CMORG-style scenario libraries the report cites — should include the prolonged unavailability of a widely used model, cloud or data provider, testing not only technical failover but the residual human capacity of many institutions to operate critical processes manually at the same time. [1, 15]

The FSB should consider a periodic implementation note or living annex — including a maintained glossary — updated through supervisory and industry evidence, rather than reopening the core principles whenever technology changes. [2, 7, 10]

Question 5. Case studies and nonbanks

Do the case studies sufficiently highlight benefits across different institutions? Are additional case studies needed, particularly for nonbanks?

Response: The case studies are useful but remain weighted toward large banks and insurers. Additional nonbank and ecosystem cases are needed.

The Bank of England/FCA survey indicates broad adoption across financial services, while IOSCO and IAIS now provide sector-specific supervisory work. The final report should use those differences to test how the sound practices operate across asset management, private markets, long-horizon asset ownership, financial market infrastructure and outsourced administration. No pension fund, sovereign wealth fund or endowment appears in the current draft, despite the report's stated application to all types of financial institutions. [7, 8, 13]

UOC recommends at least four additional case studies, and would welcome the opportunity to convene asset owners and relevant experts to help develop them:

1. **Asset management: common-model liquidity and execution risk.** Several asset managers use the same third-party model for liquidity scoring, portfolio rebalancing or execution. A model update changes risk estimates across firms and produces synchronised trading. The case should show dependency mapping, independent benchmarks, material-change notification, portfolio limits, manual fallback and a cross-market scenario exercise. [3, 7, 14]
2. **Private credit and private markets: embedded AI across the advisory chain.** A manager uses AI embedded in underwriting, valuation, covenant monitoring, fund administration and legal-document workflows. The case should address conflicts where one intermediary configures and validates the system, the rights of limited partners and boards to obtain evidence, independent review, version control, escalation and the effect of common service providers across multiple funds. [5, 7, 10]
3. **Pension fund or sovereign investor: governance purchasing power.** A long-horizon owner procures AI-enabled investment and advisory services using standard disclosure, audit, incident-notification, portability and model-change clauses. Multiple owners use pooled assurance and common vendor questionnaires. The case would demonstrate how repeat purchasers can improve market standards rather than act only as passive recipients of provider terms. [4, 5, 21]
4. **Financial market infrastructure or fund administrator: operational continuity under prolonged AI failure.** An AI-enabled process affects collateral, settlement, reconciliation, transfer agency or records. The case should test severe but plausible prolonged failure, fourth-party dependency, reconciliation controls, manual continuation against actual staffing and throughput, data portability and coordinated incident response. [4, 15]

Question 6. Actionability of case studies

Do the case studies provide actionable insights for responsible AI adoption?

Response: Partly. They illustrate practices, but a common case-study template — and methodological caveats on reported metrics — would make the insights more operational.

Each case study should state, in a consistent format:

- **Decision and stakeholders.** The decision or service being supported and the affected stakeholders.
- **Accountability map.** The accountable owner and the roles of operator, validator, approver and assurer.
- **Capability profile.** The AI capability, degree of autonomy and external systems or data it can access.
- **Materiality.** The assessment at use-case, institution, portfolio and system levels.
- **Dependency map.** Critical third-, fourth- and nth-party dependencies.
- **Preventive controls.** Pre-deployment controls, performance thresholds and conditions for suspension.
- **Evidence and escalation.** Records retained, material-change triggers and the escalation pathway.
- **Resilience and redress.** Failure scenario, fallback, exit and redress arrangements.
- **Learning.** Lessons transmitted to the board, supervisors and peer institutions.

This format would help institutions translate narrative examples into decision-rights matrices and control designs, and improve comparability across sectors and jurisdictions without turning the sound practices into prescriptive rules. [7, 8, 17, 19]

Reported metrics should carry a standing caveat. Several case studies quote striking outcome figures — detection rates above ninety-five per cent, fraud-loss reductions above twenty per cent, turnaround times cut from days to seconds. These are useful illustrations, but unaudited firm-reported metrics reproduced in an FSB document acquire a weight their provenance may not support, and risk anchoring board and supervisory expectations. A short statement that figures are self-reported and context-dependent, together with encouragement to disclose evaluation methodology alongside results, would preserve the illustrative value while protecting the report's authority. [1]

Question 7. Glossary

Are the glossary definitions clear and aligned with industry sound practices, including recent developments?

Response: The existing definitions are generally clear. Several governance and ecosystem terms should be added or refined, and terms used in the body text should be promoted to the glossary.

Terms used in the report but absent from the glossary. “Shadow AI” appears in the body text (Sections 1.3 and 3.8) and is a concept boards will need to operationalise, yet it is not defined in the glossary. “Reward hacking” and “automation bias” are defined only in footnotes despite carrying substantive weight in the agentic-risk and human-oversight discussions. The taxonomy of oversight modes in Sound Practice 10 (human-in-the-loop, human-on-the-loop, AI-in-the-loop, human-in-command) is likewise glossary-worthy, since these terms are used inconsistently across the industry. All should be promoted to the glossary. ^[1]

Recommended additions:

- **AI agent.** An AI system capable of pursuing goals through a sequence of actions, including the use of tools, data, APIs, systems or other agents, within defined or adaptive permissions.
- **Human oversight.** A control function performed by one or more competent and authorised persons who have sufficient information, independence, time and practical ability to challenge, intervene, suspend, reverse or escalate an AI-supported action.
- **Material change.** A change to a model, system, data, guardrail, prompt, tool, permission, provider, subcontractor, hosting arrangement or use case that could alter performance, risk, explainability, resilience or legal obligations.
- **Embedded AI.** AI functionality incorporated into a wider product or service, including professional, operational, investment, legal, audit, data or technology services, whether or not marketed as an AI service.
- **Systemic materiality.** The potential for an AI use or dependency to affect financial stability through scale, interconnectedness, common adoption, substitutability constraints, correlated outcomes or concentration, including where individual uses are not material in isolation.
- **Independent validation.** Effective challenge by a function sufficiently separate from the design, operation and commercial sponsorship of the AI use, with the information, competence, organisational standing and authority required to assess fitness for purpose and control effectiveness.
- **Contestability.** The ability of an affected person or institution to obtain meaningful review of an AI-supported outcome, challenge relevant evidence or assumptions, and secure correction or redress where appropriate.
- **AI incident and near miss.** An event, control failure or abnormal behaviour that caused, or could reasonably have caused, material harm, breach, disruption or unauthorised action.
- **Supervisory observability.** The capacity of relevant authorities to identify and aggregate material AI uses, dependencies, providers, incidents and concentration across institutions sufficiently to assess system-wide risk.

The definition of “supply chain” should also be reconsidered. Limiting it to services within a direct third-party relationship may be too narrow for AI ecosystems in which dependencies are indirect,

embedded, open-source, vertically integrated or introduced through subcontractors. The final definition should capture material indirect dependencies while preserving proportionality. Finally, the glossary or introduction should state expressly that “financial institution” includes asset owners such as pension funds, sovereign wealth funds and endowments. [4, 5, 7]

Proposed targeted amendments to the sound practices

The following drafting is offered to strengthen the final report without changing its principles-based character. The accompanying rationales identify both established authority and the UOC extension.

Sound Practice 2 — Governance and accountability

Financial institutions define and document clear roles, decision rights and responsibilities for AI adoption, including business ownership, operation, independent validation, approval, risk oversight and assurance. For material AI uses, approval stages appropriately separate operators, validators and approvers. Accountability remains with the financial institution and is not displaced by outsourcing or reliance on professional intermediaries, and internal accountability is aligned with the professional-conduct obligations that staff in oversight and analytical roles already carry.

Rationale. Clear accountability, independent validation and effective challenge are established elements of model-risk and AI-governance frameworks; professional codes of conduct extend that accountability to the individual level. UOC's extension is to test independence across the entire governance chain, including technology providers and advisers, rather than only within the regulated institution. [7, 9, 11, 12, 17, 19, 22]

Sound Practice 3 — Incorporation of AI risks into the risk management framework

Financial institutions maintain effective, appropriately independent channels through which employees, contractors and relevant third parties can escalate material AI incidents, near misses, undisclosed changes, control weaknesses and unresolved validation disagreements. Material themes and unresolved matters are reported to the board or an appropriate delegated committee, with protection against interference or retaliation for good-faith escalation.

Rationale. Documentation and challenge are insufficient if credible signals can be absorbed or contained before reaching accountable decision-makers. This is a UOC institutional-integrity extension, grounded in risk-governance principles, professional-conduct duties to dissociate from violations, and UOC's Charter commitment to lawful fiduciary inquiry and non-retaliation. [9, 11, 17, 21, 22]

Sound Practice 5 — Materiality and risk assessment

Financial institutions implement an effective and systematic approach to assess the materiality and risk of AI use cases at inception and thereafter, considering use-case, institution, group, portfolio and financial-system effects, including common dependencies, correlated outcomes, substitutability constraints and cumulative exposure.

Rationale. Analyses by the FSB and the Bank of England already recognise concentration and correlated behaviour, while the 2026 US interagency guidance recognizes aggregate model risk

arising from common assumptions, data and methodologies. UOC's extension is to make the transition from individual to aggregate and systemic materiality explicit in the sound practice. [2, 3, 12, 14]

Sound Practice 8 — Explainability and transparency

Financial institutions provide appropriate transparency tailored to different stakeholders, including investors and asset owners. For material AI uses, financial institutions make available information sufficient for investors to assess the institution's AI governance, materiality assessments and significant third-party dependencies, while preserving proportionality and legitimate commercial confidentiality.

Rationale. The consultation report tailors transparency to boards, customers, auditors and authorities, and notes only in passing that customer transparency “can also apply to investors”. Asset owners are a primary discipline mechanism for the governance of financial institutions; disclosure sufficient to support stewardship harnesses market discipline as a complement to supervision at no cost to proportionality. This is a UOC extension grounded in the Charter's transparency and accountability principles. [1, 16, 21]

Sound Practice 10 — Human oversight

Financial institutions implement human oversight that is demonstrably competent, informed, independent, authorised and operationally capable of intervention. For material autonomous actions, recommendation, validation, authorisation and execution are appropriately separated, with defined limits, escalation, suspension and recovery mechanisms. Where real-time review is impracticable, the institution applies constrained autonomy, pre-authorised boundaries, ex-post surveillance and automatic suspension triggers. Institutions preserve the analytical competence of staff in oversight roles, including through periodic unassisted exercise of core skills, and monitor for skill atrophy.

Rationale. International AI frameworks consistently support human oversight and accountability. UOC's extension is twofold: to define oversight as an operational capability with decision rights, resources and intervention mechanisms rather than the mere presence of a person in a workflow; and to treat the competence underpinning that capability as a control to be preserved. CFA Institute research warns that plausible-but-incorrect AI output raises the demands of genuine review and that over-reliance risks cognitive deskilling — the augmentation of human intelligence, not its substitution, is the durable model. [8, 16, 17, 18, 20, 24, 25, 26]

Sound Practice 12 — Third-party AI risk management

Financial institutions identify and manage direct and indirect risks from third-party and embedded AI, including performance, transparency, data quality, conflicts of role, material model or feature changes, supply-chain and concentration risk, incident reporting, auditability, portability, substitutability, business continuity and exit. The scope includes AI embedded within wider professional, investment, legal, audit, administrative, data, operational and technology services. Exit strategies and continuity plans that rely on manual continuation of an AI-enabled service are periodically tested against actual staffing, skills and throughput, including under scenarios of prolonged provider unavailability.

Rationale. FSB, BCBS and IOSCO work supports a broader view of third-party arrangements, supply chains and concentration, and the UK critical-third-party regime demonstrates a system-level oversight model. UOC's extension is to recognise that AI may enter a financial institution through professional and operational services even where no standalone AI system is procured. [4, 5, 6, 7, 10, 15]

Additional cross-cutting principle — System-wide dependencies and collective resilience

Financial institutions participate, proportionately, in common assurance, information-sharing, pooled testing and sector exercises where shared AI providers, models, data or infrastructure may create correlated or systemic risk. Relevant authorities should be able to aggregate material dependency information across institutions and coordinate supervisory responses to common vulnerabilities.

Rationale. Financial-stability risk may become visible only after exposures are aggregated. The FSB's monitoring work, the Bank of England's macroprudential programme and the UK CTP framework support stronger system-level visibility. UOC's extension is to link supervisory observability with the governance purchasing power of repeat purchasers and the stewardship of asset owners, so that public oversight and market discipline reinforce one another. [2, 3, 4, 14, 15, 21]

Additional cross-cutting principle — Recovery capacity and critical redundancy

For critical operations dependent on AI, financial institutions maintain defined minimum redundancies as a condition of reliance rather than as an object of cost optimisation. These include: trained human capacity sufficient to operate critical processes at defined degraded service levels; substitutable providers, models or degraded-mode arrangements; independent access to the data and records needed to reconstruct positions and decisions; and documented manual procedures that are periodically rehearsed. Redundancies are tested through realistic exercises, including scenarios of prolonged provider unavailability, with results reported to the board and reflected in continuity plans and, where appropriate, investor-facing disclosure.

Rationale. Competitive pressure guarantees that costly redundancy will be underprovided: an institution that maintains manual capacity, alternative providers and rehearsed fallbacks bears the cost individually, while the benefits of systemic resilience accrue collectively. This is the same market failure that justifies capital and liquidity buffers, which are deliberately inefficient in benign conditions. The consultation report's continuity provisions already contemplate the ability to continue AI-enabled services manually; UOC's extension is to make that capability real — tested against actual staffing, skills and throughput — and to recognise that human recovery capacity is consumed gradually as AI efficiency gains are captured, so its preservation must be deliberate. This principle complements, and is reinforced by, the professional-capacity refinements under Sound Practices 4 and 10: the skills that make oversight meaningful in normal conditions are the same skills that make recovery possible in stressed ones. [1, 4, 14, 15, 25]

Conclusion

The Consultation Report is a timely and credible foundation for responsible AI adoption in finance. Its enduring value will depend on whether it preserves institutional command as capability becomes more autonomous and more concentrated outside the regulated entity.

The final framework should therefore make six ideas unmistakable:

- **Control.** Formal accountability must be matched by practical control or reliable rights to obtain the information, authority and capability necessary to exercise it.
- **Oversight.** Human oversight must be designed as an effective decision and intervention capability, not a nominal role.
- **Professional capacity.** The competence, individual accountability and standing of the professionals who exercise oversight are controls in their own right, to be preserved as deliberately as any technical safeguard.
- **Recoverability.** The capacity to operate critical functions through prolonged AI disruption — human, procedural and technical — is consumed gradually as efficiency gains are captured, and must be maintained as deliberately as it is consumed.
- **Systemic perspective.** Financial-stability risk must be assessed across shared governance and dependency chains, not only within individual institutions.
- **Institutional learning.** Evidence, incidents and near misses must connect to independent escalation, accountable decision owners and collective learning.

With these refinements, the FSB's sound practices can do more than support safe deployment. They can establish a durable institutional architecture for preserving accountability, resilience, legitimate human judgment and public confidence as AI becomes embedded across the financial system.

UOC would welcome continued engagement with the FSB, including as a contributor to future monitoring exercises and to the development of asset-owner case studies for the final report.

Selected primary authorities

The following sources are selected for direct relevance and institutional authority. They are not intended as an exhaustive literature review.

Financial Stability Board

- [1] Financial Stability Board, Sound Practices for Responsible Adoption of Artificial Intelligence (AI): Consultation Report, 10 June 2026.
- [2] Financial Stability Board, Monitoring Adoption of Artificial Intelligence and Related Vulnerabilities in the Financial Sector, 10 October 2025.
- [3] Financial Stability Board, The Financial Stability Implications of Artificial Intelligence, 14 November 2024.
- [4] Financial Stability Board, Enhancing Third-Party Risk Management and Oversight: A Toolkit for Financial Institutions and Financial Authorities, 4 December 2023.

Financial-sector standard setters and supervisors

- [5] Basel Committee on Banking Supervision, Principles for the Sound Management of Third-Party Risk, December 2025.
- [6] Basel Committee on Banking Supervision, Digitalisation of Finance, May 2024.
- [7] International Organization of Securities Commissions, Supervisory Toolkit for AI Use in Capital Markets, Final Report, May 2026.
- [8] International Association of Insurance Supervisors, Application Paper on the Supervision of Artificial Intelligence, July 2025.
- [9] Bank for International Settlements, Consultative Group on Risk Management, Governance of AI Adoption in Central Banks, January 2025.
- [10] BIS Financial Stability Institute, Regulating AI in the Financial Sector: Recent Developments and Main Challenges, December 2024.
- [11] Prudential Regulation Authority, SS1/23 — Model Risk Management Principles for Banks, updated 16 April 2026; amendment effective 23 April 2026. The current version expressly covers externally developed and vendor models and includes independent validation as Principle 4.
- [12] Board of Governors of the Federal Reserve System, FDIC and OCC, Revised Guidance on Model Risk Management (SR 26-2), 17 April 2026. The guidance recognises aggregate model risk, effective challenge, clear roles and vendor-model validation; it excludes generative and agentic AI from its defined model scope but states that broader governance should address tools outside that scope.
- [13] Bank of England and Financial Conduct Authority, Artificial Intelligence in UK Financial Services — 2024, 21 November 2024.
- [14] Bank of England, Financial Policy Committee, Financial Stability in Focus: Artificial Intelligence in the Financial System, 9 April 2025.
- [15] Bank of England, Prudential Regulation Authority and Financial Conduct Authority, Operational Resilience: Critical Third Parties to the UK Financial Sector (PS16/24 and SS6/24), 12 November 2024. The regime provides a concrete precedent for system-level oversight of critical services without displacing firms' own accountability.

Cross-sector AI governance and legislation

- [16] Organisation for Economic Co-operation and Development, Recommendation of the Council on Artificial Intelligence (OECD AI Principles), as updated 3 May 2024.
- [17] US National Institute of Standards and Technology, Artificial Intelligence Risk Management Framework (AI RMF 1.0), January 2023.
- [18] US National Institute of Standards and Technology, Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile (NIST AI 600-1), July 2024.
- [19] International Organization for Standardization and International Electrotechnical Commission, ISO/IEC 42001:2023 — Artificial Intelligence Management Systems, December 2023.

[20] European Union, Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act), 13 June 2024.

Institutional basis

[21] Universal Owners Council Limited, Founding Charter, 18 December 2025. Articles II, III, VI and VII.

Professional standards and profession-level research

[22] CFA Institute, Code of Ethics and Standards of Professional Conduct, effective 1 January 2024.

[23] CFA Institute, Ethics and Artificial Intelligence in Investment Management: A Framework for Professionals, 2022.

[24] CFA Institute, AI in Asset Management: Tools, Applications, and Frontiers, November 2025.

[25] CFA Institute Enterprising Investor, AI in Investment Management: From Exuberance to Realism, December 2025.

[26] CFA Institute, AI in Finance: Changing Workflows, Growing Demand for Human Judgment, January 2026.