

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

Response to Consultation

Moorhouse

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?

We agree that the report accurately captures the core benefits of AI adoption, including improvements in efficiency, risk management, operational resilience and customer service. These benefits resonate strongly with how financial institutions are leveraging AI to transform business models, improve productivity and strengthen resilience. The risks highlighted, such as data bias, explainability challenges, cyber security vulnerabilities and third-party dependencies, are also well-articulated and central to our advisory work with clients. Recent industry engagement has shown AI moving beyond efficiency use cases into investment research, analytics, decision-making and client-service operations, reinforcing the need to consider both productivity benefits and new dependencies across data, platforms, oversight and operating models.

However, there are emerging risks we observed through our client work that deserve emphasis. As AI adoption matures, supervisory attention should increasingly consider systemic dependency risk alongside firm-level model risk. The growing concentration of foundation models, cloud infrastructure and AI service providers means that failures or vulnerabilities may propagate simultaneously across multiple institutions, creating new forms of operational and financial stability risk. Related to this is the potential for market correlation and herding behaviour, where reliance on common models, datasets or signals may amplify pro cyclical and reduce diversity of decision-making during periods of market stress. Chief among emerging operational risks is the rise of “Shadow AI”, unauthorised or uncontrolled adoption of AI tools by employees, which poses significant operational, conduct and data security risks within regulated environments.

There is also a “benefit dilution risk”, where some of the AI projects in our FS and adjacent regulated markets proceed without clear linkage to strategic objectives or measurable business value, leading to stalled pilots, fragmented adoption and wasted investment. Finally, the skills and cultural dimension cannot be ignored. Without structured frameworks for board literacy, second-line oversight, role-based capability development and effective challenge, firms risk underestimating both compliance gaps and operational dependencies on human judgement. In our experience updating or creating a People Framework to support the introduction and scaling of AI into BAU processes, is a structured method for

embedding AI literacy, accountabilities and developing clear escalation paths as core safeguards against these emerging risks

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

Yes. We believe the FSB's sound practices present a comprehensive and pragmatic menu for financial institutions to embed responsible AI adoption. The focus on governance, lifecycle risk management, organisational accountability and human oversight provides an appropriate high-level blueprint. However, two areas would benefit from greater clarity. This is consistent with our recent industry discussions where AI, tokenisation and infrastructure change were framed not as standalone technologies, but as operating-model transformations requiring clear accountability, control, capability and measurable value.

First: Embedding benefit-realisation checkpoints: Sound practices should go beyond risk management to require evidence of measurable value creation. Responsible adoption must not only mitigate risk but also demonstrate outcomes such as efficiency gains, capacity release, enhanced resilience and improved customer outcomes. Introducing value-realisation checkpoints at key governance stages would help firms maintain strategic alignment, avoid investment dilution and ensure AI initiatives continue to justify their risk profile.

Second: Greater guidance on proportionality and capability requirements: While the report repeatedly references proportionality, institutions would benefit from clearer guidance on applying different controls for high- versus low-materiality AI use cases. This should include expectations for role-based capability frameworks, oversight certification and Board-level AI literacy, ensuring governance remains effective as adoption scales. Institutions should also be encouraged to adopt the principle of Minimum Viable Intelligence: selecting the least complex model capable of achieving the required business outcome within acceptable cost, risk and control limits. This would reinforce proportionality, improve explainability and reduce unnecessary operational complexity.

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?

The FSB's sound practices strike an appropriate high-level balance between risks common to all forms of AI and the additional risks introduced by GenAI and agentic AI. Their focus on governance, lifecycle risk management, human oversight, cyber resilience and third-party risk provides a pragmatic foundation for responsible adoption across different AI use cases.

However, our implementation experience across financial services and adjacent regulated sectors suggests that two areas would benefit from greater emphasis. First, responsible adoption needs to be linked more explicitly to measurable value. In live AI programmes, governance can become over-weighted towards control artefacts unless it is also tied to evidence of benefit, such as improved resilience, reduced manual effort, better customer outcomes or more effective risk management. We would therefore encourage the FSB to

include value-realisation checkpoints as part of governance milestones, so that firms can evidence both responsible control and continued business justification.

Second, capability should be treated as a core control, not simply an enabling condition. In our experience supporting organisations with AI adoption, the most material implementation gaps often arise not from the absence of policy, but from unclear accountabilities, uneven AI literacy, limited second-line confidence and insufficient practical understanding among senior decision-makers. This is particularly important for GenAI and agentic AI, where risks can emerge through prompt design, tool use, workflow integration, human hand-offs and third-party dependencies. The sound practices would therefore be strengthened by more explicit guidance on role-based capability, Board and senior management literacy, and progressive upskilling for risk, technology and business teams involved in approving, deploying and overseeing AI.

These additions would not make the practices more prescriptive. Rather, they would make the existing risk-based and technology-neutral approach more executable, by recognising that responsible AI adoption depends on the interaction between governance, value, technical controls and organisational capability.

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

We believe the principles are appropriately outcome-focused and technology-neutral, providing a strong foundation that should remain relevant as AI evolves. However, flexibility should be supported by mechanisms that enable governance to evolve alongside technology. As institutions move from traditional AI to GenAI and increasingly agentic systems, governance frameworks should be reviewed whenever new AI capabilities or material operating model changes are introduced.

Our observation from AI implementation programmes across financial services and other regulated sectors is that governance frameworks can quickly become disconnected from technology reality. New capabilities are often adopted faster than governance, risk and control structures are updated, particularly where organisations move from traditional predictive models to GenAI assistants and agentic workflows.

We therefore recommend the FSB encourage periodic governance refreshes informed by structured testing, implementation experience and lessons learned from deployment. This would help ensure supervisory expectations evolve alongside practical evidence rather than static policy alone. We have also observed that institutions with stronger AI literacy, specialist expertise and executive oversight are typically better able to adapt governance as AI capabilities mature, suggesting organisational capability is a critical component of long-term flexibility.

Finally, proportionality should extend to model selection. Implementation experience suggests that organisations frequently default to more complex AI solutions than the use case requires, creating unnecessary governance, explainability and operational challenges. Institutions should therefore be encouraged to adopt the principle of Minimum Viable Intelligence: selecting the least complex model capable of achieving the intended business outcome within acceptable cost, risk and control limits. This supports explainability, reduces

unnecessary complexity, and aligns AI deployment with both risk appetite and commercial value.

- 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 included in the report outline strong examples from banking, yet coverage could be expanded to reflect adjacent sectors such as insurance and emerging hybrid models.

Case Study 1: Tier 1 Global Insurer : Agentic Claims Automation

Moorhouse supported AXA Partners in redesigning its claims process using agentic AI. At inception, straight-through processing was around 4%. Following the introduction of structured governance and a GenAI-enabled pilot, STP surpassed 12% within nine months, with a roadmap toward 30%. The programme was not limited to process efficiency—it established cross-functional accountability and rigorous value management alongside automation.

Case Study 2: Tier 1 Global Insurer: GenAI Contact Centre Augmentation

In a separate engagement, a global insurer deployed GenAI to enhance customer support during peak travel periods. A/B testing of build-versus-buy MVP solutions led to a 20% reduction in call handling times while preserving compliance and service standards.

Both cases illustrate that responsible AI adoption accelerates when governance disciplines, benefit tracking, and workforce capability are treated as design anchors. They also highlight how proportionate supervisory frameworks can actively enable innovation across nonbank financial institutions.

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

The case studies provide a useful starting point by demonstrating the relationship between governance, risk management and AI adoption. However, their value would be enhanced by including more practical examples of how financial institutions have translated responsible AI principles into operational outcomes and measurable decision-making.

In our experience supporting large financial institutions, the most actionable insights come from examples that clearly link governance interventions to business, risk and customer outcomes. For example, in a global insurance claims transformation programme, AI-enabled automation was introduced alongside structured governance, defined accountability and benefit tracking. The organisation monitored metrics such as straight-through processing rates, exception volumes, customer outcomes and model performance thresholds, enabling leadership teams to assess both the value delivered and the effectiveness of the associated controls as adoption increased.

Similarly, organisations deploying GenAI-enabled customer service capabilities have established clear operational guardrails before scaling adoption. These typically include human-in-the-loop review thresholds, accuracy and quality monitoring, compliance testing, customer complaint tracking and ongoing oversight of model outputs. By linking these controls to operational performance indicators, firms have been able to make informed decisions regarding the appropriate level of oversight, escalation and risk management required as solutions mature.

We have also observed that organisations achieve more effective outcomes when AI governance requirements are aligned to the materiality and risk profile of individual use cases. Lower-risk applications are often subject to streamlined approval and monitoring arrangements, while higher-risk customer-facing or decision-support use cases require enhanced validation, independent review and ongoing performance monitoring. Practical examples of these proportional approaches would help firms better understand how responsible AI principles can be implemented in practice.

To further increase the report's usefulness, the FSB could supplement the existing case studies with illustrative examples of:

- How organisations assess AI materiality and determine proportionate governance requirements.
- The key metrics firms use to monitor AI performance, resilience, customer outcomes and control effectiveness.
- Governance decision points that trigger additional validation, monitoring, human oversight or escalation.
- How governance, risk management and assurance models evolve as AI adoption progresses from pilot use cases to enterprise-wide deployment.

Including these implementation-focused examples would help translate high-level principles into repeatable practices, providing firms with greater confidence in designing proportionate governance and risk management frameworks that enable innovation while maintaining appropriate oversight and accountability.

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

The glossary offers a strong foundational vocabulary, and in our view the definitions provided are clear and aligned with industry sound practices. However, to ensure a clear picture is presented we suggest including the below terms in the glossary (noting some of these terms link to suggestions of additions we have made in the final question).

Suggested additions to the glossary

Production architecture and operational risk

-AgentOps: The set of practices, processes, and tools used to develop, deploy, monitor, evaluate, and improve AI agents in production. AgentOps focuses on reliability,

accountability, observability, and control of agent behaviour over time, especially where agents act autonomously, use tools, or interact with external systems.

-LLMOps: The set of practices used to manage large language model applications across their lifecycle, including model selection, prompt management, testing, deployment, monitoring, security, versioning, and change control. LLMOps extends conventional MLOps to address the stochastic behaviour, prompt sensitivity, and third-party dependencies of LLM-based systems.

-Evaluation framework (Evals): A structured approach for testing whether an AI model, system, or agent performs acceptably for its intended use. Evals may include benchmark tests, task-based assessments, expert review, red-teaming, comparison against alternative models, and repeated testing over time to assess quality, reliability, safety, and suitability for deployment.

-Observability: The ability to inspect, trace, and understand how an AI system behaves in operation through logs, metrics, traces, records of inputs and outputs, tool-use history, intermediate steps, and alerts. Observability supports monitoring, incident investigation, root-cause analysis, auditability, and effective human oversight.

-Tool use: The ability of an AI model or agent to call external functions, APIs, databases, software tools, or digital systems to retrieve information, take actions, or complete a task. Tool use can improve capability and utility, but also expands operational, cyber, and conduct risk.

-Memory: A mechanism through which an AI system stores, retrieves, and reuses information across steps or interactions. Memory may include short-term task context, longer-term user or workflow context, and retrieved external knowledge. Memory can improve continuity and performance but may also introduce privacy, security, and poisoning risks.

-Planning loop: An iterative process in which an AI agent decomposes a goal into steps, executes actions, reviews results, and revises its plan until a stopping condition is met. Planning loops can improve problem-solving but may create latency, cost, unpredictability, and control challenges if poorly bounded.

-Orchestration: The coordination layer that manages how models, prompts, tools, memory, workflows, human checkpoints, and other agents interact to complete a task. Orchestration determines sequencing, routing, fallback behaviour, and control logic across an AI system.

-Multi-agent system: A system in which two or more AI agents interact, coordinate, or divide labour to achieve a shared or related objective. Multi-agent systems may improve flexibility and scale, but can introduce added complexity, debugging difficulty, and unforeseen interactions or failure modes.

-Model Context Protocol (MCP): A standardised interface for connecting AI models or agents to external tools, data sources, and services through a consistent interaction pattern. MCP aims to reduce integration complexity and improve portability, governance, and control over how models access context and capabilities.

Economics, cost, and governance

-Token consumption: The volume of tokens processed by an AI model, including input, output, and, where applicable, hidden reasoning or multi-step agent activity. Token consumption is a key driver of cost, latency, and scalability in LLM-based systems.

-Model selection: The process of choosing the most appropriate model or system for a use case by balancing performance, explainability, cost, latency, risk, operational requirements, and availability. Model selection may also consider whether simpler or lower-cost models are sufficient for the desired outcome.

-Cost management: The practices used to measure, control, and optimise the commercial, computing, energy, staffing, and third-party costs of AI systems through design choices, usage controls, monitoring, and governance.

-AI FinOps: The discipline of managing the financial performance of AI systems in production, including spend visibility, usage governance, cost allocation, efficiency optimisation, and alignment of cost to business outcomes. AI FinOps is particularly relevant for LLM and agentic systems where token usage, latency, retries, and tool calls can create rapidly escalating operating costs.

-Cost-per-outcome: A measure of AI efficiency that relates total operating cost to a valid business result, such as a resolved customer query, a reviewed claim, a completed underwriting memo, or a fraud rule produced and approved. Cost-per-outcome is often more decision-useful than model unit cost alone because it captures the full economics of delivering value.

-Minimum Viable Intelligence (MVI): The minimum level of model capability, autonomy, and system complexity required to achieve a use case's target outcome within acceptable cost, risk, and control limits. MVI encourages firms to avoid over-specifying model sophistication where simpler approaches are sufficient.

Organisation, adoption, and operating model

-AI literacy: The baseline understanding staff need in order to use AI appropriately in line with relevant compliance, recognise its limitations, identify risks, and escalate concerns. AI literacy includes knowing when AI may be helpful, when outputs require verification, and how governance boundaries apply.

-Digital fluency: The practical ability to work confidently with digital tools, data, and AI-enabled workflows in day-to-day work. Digital fluency is broader than awareness; it includes applying tools effectively, understanding process implications, and adapting behaviours as work changes.

-Human-AI teaming: A mode of work in which humans and AI systems complement one another through clear division of roles, hand-offs, escalation paths, and feedback loops. Effective human-AI teaming depends on task design, incentives, authority to intervene, and the ability of humans to challenge or override AI outputs when needed.

-AI operating model: The organisational structure, decision rights, governance mechanisms, delivery model, and accountabilities through which an institution develops, deploys, oversees, and scales AI. An AI operating model may be centralised, decentralised, or hybrid, and should align with strategy, risk appetite, and the scale of AI adoption.

-Adoption: The process by which AI moves from experimentation to sustained use in real work, including changes to workflows, incentives, controls, user behaviour, and management routines. Adoption is not simply technical deployment; it requires that users integrate the system into how work is done. This should be considered from the perspective of both teams and people.

-Change management: The structured effort required to help people, teams, and processes adapt to AI-enabled ways of working. This includes communication, training, role redesign, managerial reinforcement, leadership endorsement and feedback mechanisms to ensure responsible and effective use at scale.

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

We welcome the FSB's continued momentum in shaping a principles-based, technology-neutral framework, and we see a strong opportunity to work together to bring these ideas into practical application. Looking ahead, we believe there is real value in developing operational toolkits that help bridge the gap between high-level principles and effective on-site assurance.

In our view, three enablers could make a meaningful difference:

A role-based capability map, grounded in governance and control responsibilities, to help bring Sound Practice 4 on organisational adaptability to life in a clear and practical way.

Configurable templates that link AI materiality tiers to risk gating structures, enabling teams to move confidently from conceptual guardrails to clear, executable decision protocols.

Exploration of shared industry utilities for model validation and explainability, supporting a more consistent approach while reducing duplication and easing compliance costs across jurisdictions.

We also value the FSB's leadership in convening cross-border sandbox initiatives. There is an exciting opportunity to broaden their scope to include GenAI and agentic systems. By doing so, these environments can become powerful spaces for collaboratively testing resilience scenarios for emerging AI agents before any systemic exposure.

We encourage the FSB to consider convening an international Responsible AI Sandbox Network. National initiatives are already demonstrating the value of structured experimentation, but many of the emerging risks associated with GenAI and agentic systems are inherently cross-border. A coordinated network would enable supervisors and firms to evaluate common scenarios, develop shared testing methodologies, exchange evidence on emerging risks, and generate supervisory learning before these technologies reach systemic scale. Such a capability would complement the sound practices by providing an evidence-based mechanism for their continuous evolution.

Together, this approach can strengthen prudential outcomes while supporting innovation, helping organisations move forward with confidence and clarity.

People side of AI/tech transformation

Based on the FSB consultation report's focus on governance, risk management, controls and AI lifecycle oversight, one area that would benefit from greater emphasis is the people and organisational dimension of AI adoption. While the report appropriately recognises the importance of governance and accountability, it largely treats AI implementation as a technical and risk-management challenge rather than a broader organisational transformation. In practice, the successful and responsible adoption of AI depends not only on policies, controls and monitoring, but also on whether people understand, trust and are able to use AI effectively in their day-to-day work.

Research and industry experience consistently show that barriers to AI adoption are often cultural, behavioural and skills-related rather than technical. Financial institutions therefore need to consider capabilities such as AI literacy, digital fluency, communication, training, leadership enablement, role redesign, resistance management and the development of organisational cultures that encourage responsible experimentation and challenge. These themes are central to enabling AI to move from isolated use cases to sustainable adoption at scale.

We would therefore encourage the FSB to strengthen the report by explicitly recognising the role of adoption, change management and operating model design as enablers of responsible AI adoption. This could include incorporating concepts such as AI literacy, digital fluency, AI operating models, adoption and change management into the glossary and relevant sound practices. Doing so would help institutions address not only whether AI systems are governed appropriately, but also whether the workforce has the skills, confidence, understanding and support required to use them safely and effectively.

Given that human judgement remains a critical control within many financial services processes, greater consideration of the people side of AI would strengthen the report's objective of promoting responsible AI adoption and long-term financial stability. 'Human'-related areas such as leadership, AI and digital literacy, culture and mindset, skills development, communication, champion networks and resistance management are core components of successful AI adoption, these concepts currently appear underrepresented in this report.