

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

Response to Consultation

The Personal Investment Management & Financial Advice Association (PIMFA)

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?

PIMFA broadly agrees with the benefits and risks identified in the report and welcomes its balanced assessment of the opportunities and challenges associated with AI adoption in financial services.

There are, however, a number of areas that could be given greater emphasis. First, the report could more clearly distinguish between different categories of AI technologies, recognising that different architectures may present different governance, explainability and risk management considerations. Second, greater consideration could be given to the growing dependence on a relatively small number of third-party AI providers and the operational challenges that may arise where providers alter models, functionality or access arrangements. Third, the report could further explore the interaction between regulated firms and AI-enabled activities that sit outside the current regulatory perimeter, including consumer use of publicly available AI tools and uncertainty around the boundaries between information, guidance and advice. The report could also recognise the importance of maintaining a level playing field as AI develops, ensuring that comparable activities with comparable risks are subject to consistently applied regulatory expectations.

The report could also recognise the importance of AI literacy and capability-building within firms, as well as the risks this poses for human capability and consumer agency more broadly. As AI becomes embedded in financial decision-making, there is a risk that both consumers and staff become less able to understand, challenge or recover from AI-driven decisions. This is particularly relevant where agentic finance is used in high-stakes areas such as pensions, investments, debt advice and customer support. We would like to highlight that effective oversight is likely to depend not only on clear accountability structures but also on ensuring decision makers have sufficient understanding to exercise informed judgement and challenge.

A related risk is the potential erosion of operational knowledge within firms. If critical processes become heavily AI-enabled, firms may retain technological resilience but lose practical human capability to operate, scrutinise or recover those processes when AI outputs

become unreliable or unavailable. This is a resilience issue for firms as much as it is a training issue.

The report may also benefit from more explicit recognition of ecosystem-level risks arising from agent-to-agent interaction. Future risks may not arise solely from individual AI models but from networks of agents interacting with firms, consumers, infrastructure providers, and each other in ways that create emergent behaviours, the rapid propagation of errors, or the concentration of market influence.

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

Overall, the sound practices provide a useful and comprehensive framework for responsible AI adoption.

We particularly welcome the principles-based approach, which provides flexibility while recognising the importance of governance, risk management and accountability. We would suggest, however, that AI literacy, competence management and challenge capability be treated as core controls. In our view, governance structures alone may not be sufficient to ensure responsible adoption unless board members, senior managers and operational users have sufficient understanding to identify limitations, challenge outputs and recognise when automation is inappropriate. The report would benefit from additional guidance in a small number of areas, including accountability across complex AI supply chains, explainability expectations for different use cases, and the governance implications of significant changes made by third-party providers.

Further examples of how proportionality should be applied in practice would also be helpful, particularly for smaller regulated firms. This should also recognise differences in firms' resources, technical capability and access to specialist expertise.

We would also suggest that the FSB consider a clearer expectation for firms to maintain an AI authority or decision-rights framework. This would help firms define which activities:

- may be supported by AI,
- may be delegated to AI subject to approval; and
- must remain human-led.

The framework could also more explicitly require AI exit, rollback and degraded-operation planning for material AI use cases. This would include the ability to operate critical processes without AI, maintain manual fallback arrangements, and test whether human expertise remains sufficient to take over when required.

Finally, we would encourage greater emphasis on AI incidents and, crucially, near-miss reporting. Firms should be encouraged to capture leading indicators such as hallucinations detected before release, inappropriate agent actions, prompt injection attempts, model drift, unexplained performance degradation and human challenge failures.

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 balance is broadly appropriate.

However, the report could place greater emphasis on recognising that AI is not a single technology and that different forms of AI may present different opportunities and risks. Governance expectations should remain proportionate to both the use case and the characteristics of the technology being deployed.

We support the report's increased focus on emerging technologies such as generative and agentic AI, while recognising that many firms continue to rely on more established AI and machine learning techniques that may warrant different governance considerations.

In developing this distinction further, we would suggest that GenAI and agentic AI could be framed using an autonomy spectrum, similar to that used in the UK FCA Mills Review (<https://www.fca.org.uk/publication/corporate/the-mills-review.pdf>). This spectrum ranges from AI as a tool used by a human operator (most often the case with GenAI), through collaborative and advisory models, to AI acting subject to human approval, and ultimately to AI operating within defined boundaries while humans monitor outcomes (more typical of agentic AI).

In our view, this is worth considering as the nature of risk appears to change with increasing autonomy. At lower levels of autonomy, risk is often linked to output quality and user over-reliance. At higher levels, risks increasingly relate to delegated authority, boundary-setting, escalation, auditability, reversibility and whether human oversight remains meaningful in practice.

We would welcome further clarity on how firms should determine appropriate levels of autonomy by use case, including consideration of customer impact, reversibility, materiality, explainability and the ability of humans to intervene effectively.

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

The principles-based approach provides an appropriate level of flexibility and should remain adaptable as technologies continue to evolve.

One area that may benefit from further consideration is the ability of firms to identify and respond to significant changes made by third-party providers. As AI models and services evolve rapidly, firms may need to move away from periodic reassessment towards a model of continuous reassessment, revisiting governance, monitoring and approval decisions whenever material model updates, provider changes, new capabilities or changes in system behaviour occur.

The framework could also recognise that AI can remain available whilst still becoming unreliable. Resilience testing should therefore include AI withdrawal and AI degradation scenarios, including model drift, model creep, altered reasoning behaviour, unexpected agent behaviour and external provider changes.

For agentic systems, firms may need stronger monitoring of behaviour over time, including whether agents remain within intended boundaries, how they escalate decisions, and whether cumulative interactions between multi-agent systems create risks that were not visible at deployment.

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 and provide practical illustrations of how responsible AI adoption can be implemented.

However, they are primarily focused on larger financial institutions. The report would benefit from the inclusion of additional examples from wealth managers, financial advisers, platforms and other non-bank financial institutions, such as AI-enabled client triage, suitability support, customer communications review, transfer processing, financial crime alert triage and operational exception handling. Case studies demonstrating proportionate governance arrangements for smaller firms would be particularly valuable.

Examples involving customer-facing applications, decision support tools and firms operating within highly regulated consumer environments would also strengthen the report.

A further useful addition would be a case study on proportionate agentic AI in consumer journeys, for example, how a firm interacts with a consumer-appointed AI agent, manages information boundaries, assesses advice or guidance thresholds, and maintains consent, accountability and audit trails.

Finally, a case study on AI withdrawal or degradation would be valuable: how a smaller firm identifies that a tool has become unreliable, moves to a manual or reduced-capacity model, preserves customer outcomes, and records lessons learned.

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

The case studies provide useful examples of governance structures, oversight arrangements and implementation considerations. We particularly welcome the focus on practical application rather than theoretical principles.

Additional examples covering smaller firms and a broader range of business models would further enhance their usefulness and help demonstrate how proportionality can operate in practice.

The case studies would also become more actionable if they included concrete implementation artefacts, such as proportionality assessments, decision-rights matrices, human oversight models, near-miss logs, model-change approval processes and fallback playbooks.

We would also welcome practical examples of what effective human oversight looks like in different contexts, including what information the human reviewer receives, how challenge

is recorded, when escalation is required, and how firms can avoid oversight becoming a tick-box exercise in which users approve AI recommendations.

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

The glossary is generally clear and accessible.

However, the report may benefit from greater precision when describing different categories of AI technologies. Recognising distinctions between deterministic systems, statistical models and newer generative or agentic technologies may help support more proportionate governance and risk management approaches. This could be achieved by defining terms that are likely to become increasingly important in practice, including agentic AI, AI autonomy, digital worker, human oversight, human-in-the-loop, human-on-the-loop, model drift, model creep, AI near miss, AI incident, AI withdrawal scenario and AI decision rights. It would also be helpful to define AI literacy and AI competence, given the importance of human understanding in delivering effective governance and challenge. Without clearer terminology, there is a risk firms may interpret these expectations inconsistently.

Given the rapid pace of technological change, maintaining terminology that is sufficiently flexible while remaining technically accurate will be important.

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

PIMFA welcomes the FSB's work on responsible AI adoption and supports the overall direction of the report. The principles-based approach provides an appropriate foundation for firms seeking to adopt AI responsibly while remaining adaptable to future developments.

As AI adoption continues to accelerate, we encourage continued focus on proportionality, recognising the diversity of firms operating within financial services. Governance arrangements appropriate for large internationally active institutions may not always be proportionate for smaller regulated firms, including wealth managers, advisers and specialist financial services providers.

We also encourage further consideration of accountability across the AI supply chain, the implications of third-party dependencies, the interaction between AI and activities outside the current regulatory perimeter, and the role of capability building in supporting effective governance.

There appears to be an emerging theme around balance: supporting competition, innovation, consumer access and operational efficiency, while preserving market integrity, customer understanding, resilience and accountability. The FSB's framework provides a strong foundation, but we believe responsible adoption will increasingly need to extend beyond models and providers to include human capability, decision rights, autonomy levels, agent ecosystems and operational knowledge resilience. We believe human factors are likely to be at least as significant a point of failure in responsible AI adoption as technology itself, and the final report would benefit from addressing AI adoption more explicitly through this human lens.

We would therefore encourage the FSB to strengthen the final report by recognising AI literacy as a core control, encouraging capability-retention and AI-withdrawal testing, expanding its treatment of agentic ecosystem risks, and providing firms with practical guidance on decision rights, incident learning and human-digital hybrid workforce governance.

AI has the potential to improve efficiency, resilience, customer service and access to financial guidance when deployed appropriately. The objective should be to support innovation through practical, proportionate and outcomes-focused governance that enables firms to realise these benefits while maintaining trust, accountability and market integrity.

PIMFA welcomes the FSB's work on responsible AI adoption and supports the report's overall direction. The principles-based approach provides an appropriate foundation for firms seeking to adopt AI responsibly while remaining adaptable to future developments.

As AI adoption continues to accelerate, we encourage continued focus on proportionality, recognising the diversity of firms operating within financial services. Governance arrangements appropriate for large internationally active institutions may not always be proportionate for smaller regulated firms, including wealth managers, advisers and specialist financial services providers.

We also encourage further consideration of accountability across the AI supply chain, the implications of third-party dependencies, the interaction between AI and activities outside the current regulatory perimeter, and the role of capability building in supporting effective governance.

There appears to be an emerging theme around balance: supporting competition, innovation, consumer access and operational efficiency, while preserving market integrity, customer understanding, resilience and accountability. The FSB's framework provides a strong foundation, but we believe responsible adoption will increasingly need to extend beyond models and providers to include human capability, decision rights, autonomy levels, agent ecosystems and operational knowledge resilience. Some of our members consider that human factors are likely to be at least as significant a point of failure in responsible AI adoption as technology itself, and the final report would benefit from addressing AI adoption more explicitly through this human lens.

We would therefore encourage the FSB to strengthen the final report by recognising AI literacy as a core control, encouraging capability-retention and AI-withdrawal testing, expanding its treatment of agentic ecosystem risks, and providing firms with practical guidance on decision rights, incident learning and human-digital hybrid workforce governance.

AI has the potential to improve efficiency, resilience, customer service and access to financial guidance when deployed appropriately. The objective should be to support innovation through practical, proportionate and outcomes focused governance that enables firms to realise these benefits while maintaining trust, accountability and market integrity.