



FINANCIAL
STABILITY
BOARD

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

Response to Consultation

Manulife

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?**
2. **Are the sound practices sufficiently comprehensive and clear to enable financial institutions' responsible AI adoption?**
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?**
4. **Are the sound practices sufficiently flexible to accommodate and address newer types of AI and responsible adoption over time?**
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.**
6. **Do the case studies in this report provide actionable insights for financial institutions in their responsible AI adoption?**
7. **Are the definitions in the glossary clear and aligned with industry sound practices, including recent developments in AI?**
8. **Are there any comments you would like to make on this report? Please fill in.**

Manulife supports the FSB's efforts to advance responsible AI adoption in financial services and welcomes the paper's balanced framing of innovation and risk management as complementary objectives. Many of the report's themes align closely with the governance, oversight, and risk-based practices Manulife has been building as AI adoption expands across underwriting, claims, distribution, customer service, investment management, and software engineering.

One of the strongest aspects of the FSB paper is its emphasis on proportionality. Not every AI use case presents the same risk profile, and there is a significant difference between a

colleague productivity tool that supports internal drafting and an AI use case involved in underwriting, claims decisions, or customer outcomes.

As supervisory and industry practices mature, it will be important to ensure that well-intentioned governance does not inadvertently create friction for lower-risk use cases without meaningfully improving risk outcomes. A proportional approach is critical to enabling innovation while maintaining appropriate controls and accountability.

Manulife strongly agrees with the report's emphasis on data governance. In many respects, responsible AI starts well before model development, because the quality, stewardship, access, and control of data directly affect the quality and trustworthiness of AI outcomes.

Foundational investments in data governance are therefore as important as the sophistication of the model itself. For organizations moving from experimentation to scaled adoption, data governance is not only an enabler of responsible AI; it is a prerequisite for sustainable and trusted implementation.

Financial institutions and supervisors need a practical approach to maintaining human agency and accountability as AI portfolios scale from dozens of use cases to hundreds or thousands that create an extremely wide range of risk exposures. Given the powerful and super fast-evolving nature of AI, the key question becomes how oversight is designed and evidenced effectively.

Meaningful oversight can take many forms, including governance, testing, escalation protocols, exception management, and ongoing monitoring. It does not always require a human to review every AI output, particularly where controls, escalation pathways, and accountability mechanisms are appropriately designed.

Manulife recognizes the importance of explainability. With generative AI and agentic AI, however, explainability is rarely binary, and some valuable AI capabilities may not be inherently explainable in the same way as more traditional models have been because of their probabilistic nature.

Even where individual outputs are challenging to explain fully at a granular mechanical way, trust can be supported through strong governance, validation, monitoring, and testing of outcomes. We encourage the FSB to continue recognizing the broader control environment as an important path to explainability, particularly as interpretability and alignment science continue to develop.

The report appropriately identifies third-party risk as one of the most significant challenges facing financial institutions. Many firms are building AI capabilities on top of technologies provided by cloud providers, providers of foundation models, AI platforms, and other software vendors, often combining third-party capabilities with internally developed intellectual property and controls.

This environment requires diligent and ongoing third-party risk management, particularly given the pace of innovation and the concentration of critical AI capabilities among a relatively small number of providers. While financial institutions remain accountable for outcomes, their ability to influence the design, governance, transparency, resilience, and

risk management practices of these technology providers and to monitor the developments in these areas on an ongoing basis may be limited, particularly where providers serve global markets and set standards at scale.

For this reason, we encourage the FSB to consider the practical constraints institutions face when managing risks that originate within third-party AI ecosystems. In addition to promoting stronger third-party risk management practices within financial institutions, there may be value in exploring mechanisms to increase transparency, standardize disclosures and assurance reporting, and, where appropriate, strengthen expectations for technology providers that deliver critical AI services to the financial sector. Greater alignment across financial institutions, regulators, and technology providers could help reduce duplication while improving transparency, consistency, and risk outcomes across the industry. Most importantly, it would contribute to the overall stability of the financial system.

Manulife sees significant potential in agentic AI as financial institutions begin deploying systems that can coordinate multi-step workflows, invoke tools, interact with operational processes, and act with increasing independence. These capabilities can create meaningful benefits, but they also raise questions about safety, operational resilience, accountability, and sustainable deployment.

As adoption matures, governance challenges will extend beyond individual agents to increasingly interconnected ecosystems of agents operating across platforms, organizations, and third-party providers. We would welcome broad industry discussions and ultimately further FSB guidance on questions such as: How should accountability be allocated when multiple agents contribute to a single outcome? What governance expectations should apply when agents interact with or delegate tasks to third-party agents? How should firms assess and manage risks arising from emergent behaviors, dynamic decision pathways, and interconnected agent networks? And what standards for interoperability, authentication, authorization, and operational resilience may be necessary as agents increasingly interact across institutional and technological boundaries?

Greater clarity on these issues would help institutions establish governance frameworks that remain effective as AI systems become more autonomous, interconnected, and capable of coordinating complex activities with limited human intervention.

The FSB paper provides a strong foundation by grounding the discussion in practical governance and highlighting data governance, proportionality, explainability, third-party risk, and emerging autonomous systems. As AI adoption accelerates, further discussion would be valuable on how sound practices can be implemented consistently at enterprise scale.

In particular, the industry would benefit from further discussions on operating models that can support thousands of AI use cases or large numbers of AI agents, reasonable oversight resourcing, which controls may be acceptable to automate, when AI-overseeing-AI may be appropriate, and what evidence firms should provide to supervisors to demonstrate that oversight remains effective.

Providing practical examples would promote greater consistency in implementation while preserving flexibility for institutions of varying sizes, business models, and AI maturity levels.

This would be particularly valuable for emerging technologies such as generative and agentic AI, where governance practices continue to evolve rapidly.

Manulife supports the FSB's objective of promoting responsible AI adoption while enabling financial institutions to continue innovating for the benefit of customers, colleagues, advisors, and markets. The principles in the paper are clear; the next challenge is ensuring they can be implemented in a practical, scalable, and proportionate manner across institutions with diverse AI portfolios.

AI is evolving quickly, and firms must keep pace with innovation while maintaining appropriate governance and risk management. Continued dialogue between industry and regulators will be essential to support responsible adoption, supervisory clarity, and effective risk outcomes as AI capabilities continue to evolve. Manulife is committed to such dialogue and will gladly contribute to it as we all work towards constructive and sustainable outcomes pertaining to the adoption of AI.