

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

Response to Consultation

Maria Luz M (independent response)

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?

The benefits and risks described are broadly accurate, but the report's risk taxonomy is missing a distinct category: the risk that a fully compliant, binding governance framework still fails because the monitoring behind it isn't operationally real. This is not hypothetical. In 2020, an Australian bank operating under existing, binding anti-money-laundering law, with real penalties attached, was fined AU\$1.3 billion (then the largest civil penalty in Australian corporate history) after admitting to 23 million statutory breaches, including a known child-exploitation risk indicator on one payment platform that went unmonitored for five years. The law existed; the enforcement infrastructure behind it did not. I recommend the FSB name this explicitly as a “governance-monitoring gap”: a category distinct from “inadequate governance,” describing cases where accountability is assigned on paper but not verified in operation.

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

Sound Practice 2 stops at organisational role clarity (“clear roles and responsibilities”) but does not require named individual accountability. Several jurisdictions already go further: Australia's Financial Accountability Regime, itself a direct legislative response to an earlier accountability failure identified by the 2018–19 Royal Commission into Misconduct in the Banking, Superannuation and Financial Services Industry, requires regulated entities to register specific accountable persons against specific responsibilities, so that a governance failure traces to an identifiable individual rather than the organisation collectively. I recommend SP2 explicitly distinguish organisational role clarity from named-person accountability, and reference existing named-accountability regimes (e.g., Australia's FAR, the UK's Senior Managers and Certification Regime) as implementation models institutions and supervisors can draw on.

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 report rightly acknowledges that continuous human monitoring of individual agentic decisions becomes impractical at scale, and its suggestion of AI-monitoring-AI architectures is appropriate. However, the sound practices do not require institutions to verify that oversight mechanisms are actually functioning, as distinct from merely existing on paper, the same failure mode I describe in my response to Question 1. I recommend SP9 (performance management) be explicitly linked to SP2/3 (accountability), so that a monitoring or performance failure is treated as an accountability event traceable to a named responsible person, not merely logged as a technical issue.

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

The principles-based, outcomes-oriented design is the right choice, and I would caution strongly against any future move toward more prescriptive, technology-specific rules. Prescriptive regimes tied to specific technical thresholds or architectures are structurally vulnerable to what the regulatory literature terms the institutional pacing problem: the gap between multi-year legislative cycles and technology that changes on a scale of months. A principles-based standard tied to institutional outcomes, accountability, monitoring, tolerance-setting, rather than to a specific technology's characteristics ages far better, provided it is periodically tested against new developments. Australia's own prudential regulator offers a working precedent for this: APRA conducted a targeted post-implementation review of its own operational risk standard within roughly two years of commencement, resulting in specific amendments based on what had been learned in practice. I recommend the FSB commit to a similarly defined periodic review cycle, for example, biennial, for the sound practices themselves, rather than leaving revision timing open-ended.

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.

Every case study in the report illustrates a successful implementation. None examines a governance or monitoring failure and the corrective response that followed, yet failure analysis is often more instructive for institutions building their own controls than replicated success stories alone. I recommend the FSB include at least one anonymised case study built around a supervisory or enforcement action, showing what failed and what changed afterward. For a nonbank example, I would suggest an insurance-sector case study: insurers already operate AI-driven underwriting and pricing models under existing anti-discrimination law requiring genuine actuarial justification for any differentiated outcome, illustrating a sector where responsible AI adoption is being enforced through pre-existing general law rather than AI-specific rules, a model directly relevant to the report's broader approach.

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

The existing case studies are informative on implementation mechanics but, per my response to Question 5, would be more actionable if paired with failure-mode analysis. An

institution reading only success stories has limited ability to pressure-test its own controls against realistic failure scenarios before they occur.

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

I recommend the glossary more sharply distinguish “agentic AI” from “generative AI” and conventional automated decision systems using an operationally meaningful criterion, degree of autonomous action taken without human sign-off, rather than capability alone, since several sound practices (particularly SP9 and SP10) apply with different force depending on this distinction. I would also recommend an explicit definition of “human oversight” itself, given that its practical meaning is precisely what is contested as agentic systems scale, per my response to Question 3.

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

The consultation report’s core premise, that responsible AI adoption should be governed through existing risk-management principles rather than new AI-specific statutes, is sound, and consistent with the approach several FSB member jurisdictions, including Australia, have independently taken. My central recommendation across this submission is that the report’s next iteration place more weight on the distinction between an institution having a sound practice and operationalising it. Governance frameworks that exist in policy documents but are not independently tested in operation have, in at least one documented case, coexisted with billions of dollars in undetected statutory breaches over several years. I recommend the final report include an explicit expectation that institutions periodically and independently test, rather than merely document, that their AI governance and monitoring controls function as designed.