

Questions	Comments
1	We generally agree with the content of the report. That said, we would like to point out the following: We believe it would be helpful to provide a clearer breakdown of the "the implications of not adopting AI" mentioned in the Introduction (page 3). For example, there are areas—such as fraud detection measures—where the mere failure to utilize AI could itself pose a risk of not fully meeting supervisory expectations. If the report were to present a framework for weighing the 'risks of not adopting AI' against the 'risks associated with adoption', it would serve as an even more valuable reference for financial institutions in making responsible decisions regarding AI adoption. Furthermore, there is a risk that AI-generated output could infringe on the intellectual property rights of other companies.
2	Regarding Sound Practice 1 (Strategic direction and oversight), in the first paragraph on page 19 titled "The board and senior management consider AI risks when setting risk appetite and tolerance", "automated decision-making in critical business areas" is cited as an example of a prohibited case of AI usage. We would like the FSB to clarify that this is merely an example and that decisions should be made based on factors such as materiality, customer impact, regulation, and compensating controls. The lack of clarity regarding what is prohibited could potentially hinder the practical use of AI. While it is important for the board and senior management to be involved in establishing frameworks such as AI utilization policies, risk appetite, oversight of key use cases, and escalation protocols, we agree with the intent of the draft. That said, to clarify practical understanding, we believe it would be helpful to specify that the involvement of the board and senior management does not uniformly require detailed judgment on individual usage cases, but rather is premised on allowing business units and control functions to make appropriate judgments and manage matters based on factors such as materiality, risk, and customer impact.
3	Regarding Sound Practices 3 (Integration of AI Risks into Risk Management Frameworks), on page 23, the statement, "Where such controls are managed by third-party vendors and rollback or configuration access is limited, financial institutions establish contractual arrangements for vendors to provide timely information and compensating controls.", when using third-party AI providers—particularly large-scale providers offering foundational models or cloud services—it is important for financial institutions to ensure, to the extent possible, that contractual provisions for information disclosure, change notifications, and configuration management are in place. On the other hand, it is conceivable that, due to market practices and other factors, it may not be practically feasible to fully secure the necessary information and response measures. Therefore, it would be useful if the report presented an approach where financial institutions mitigate risk by combining measures such as limiting the scope of use, verifying output, conducting continuous monitoring,

	establishing reporting and response processes in the event of issues, and setting conditions for suspending or reviewing usage as necessary. Regarding Sound Practice 10 (Human oversight), this report outlines forms of supervision for agent AI, such as "human-in-command", while also acknowledging that as the use of agents expands, real-time human oversight may become impractical. In light of this, regarding the transition from human oversight to AI-based oversight ("AI-in-the-loop"), since AI-based monitoring itself poses challenges such as model risk, explainability, and the assignment of responsibility, we believe it would be useful to provide practical examples—including the independence of monitoring AI, performance verification, log recording, and escalation to humans—while ensuring that ultimate accountability remains with humans and organizations. Furthermore, it would be highly beneficial for financial institutions utilizing agents if the industry and regulatory authorities were to provide in-depth and concrete guidance on building consensus regarding societal acceptance of the "human-in-the-loop" framework. Page 28 describes limiting AI's access to data as a risk mitigation measure, and page 54 lists measures such as limiting who can use AI (e.g., those who have undergone training). However, since restricting access could reduce the effectiveness of AI, it is desirable to implement tiered controls based on data classification and the use cases to avoid imposing excessive restrictions even on low-risk applications aimed at improving productivity. Visualizing AI usage cases and managing an AI inventory are considered useful tools that contribute to the identification, assessment, monitoring, and management of AI risks. On the other hand, requiring the same level of documentation and lifecycle management for all AI usage could impose an excessive practical burden, particularly for low-risk productivity-enhancing applications and limited internal use. Therefore, it would be helpful to clarify that the scope and granularity of AI inventories and documentation can be flexibly determined based on factors such as materiality, risk, impact on customers, and dependence on third parties.
4	Although this report is positioned as a set of sound practices designed to ensure flexibility, there is concern that the detailed provisions (particularly Sound Practice 10 regarding human oversight) may be used by supervisory authorities as de facto compliance checklists. If this occurs, it could hinder the use of AI, undermining the framework's intended flexibility. We request that it be explicitly stated that these sound practices 'should not be used as comprehensive compliance checklists'. Regarding the eighth bullet point listed under "Additional considerations for GenAI and agentic AI" on p. 43, based on the explanation of "Human-in-the-loop" on p. 42, it is considered that this may also apply to the handling of insurance claims. Given the evolution of AI, it is conceivable that, depending on the nature of the insurance policy, human approval or dual approval may no longer be required in the future.

	Rather than always requiring human approval or dual approval, we suggest explicitly stating that for low-risk, small-amount, or routine insurance claims, post-hoc monitoring, sampling, and manual escalation when thresholds are exceeded may be permitted in lieu of prior approval. We would like it to be clarified that the specific controls regarding GenAI and Agentic AI are merely examples, and that in practice, each company may adopt controls appropriate to its own materiality and risk profile.
5	Although case studies are useful, the examples presented tend to be biased toward success stories with quantifiable results, and there are only a limited number of cases involving decisions to postpone or halt AI adoption. From the perspective of responsible implementation, sharing cases of withdrawing or not adopting AI would likely provide valuable practical insights. At the same time, since presenting cases of withdrawal or non-adoption could be mistakenly interpreted as future prohibitions or restrictions, it is important to clearly convey—without any misunderstanding—that the intent is not to prohibit specific technologies or usage cases. Therefore, rather than simply stating the conclusion that "the use case or technology type in question is inappropriate", please clarify that such decisions are context-dependent by specifying the preconditions at the time of the decision, the perspectives of risk assessment, the availability of alternative controls, constraints related to data, models, and business processes, and the possibility of reconsideration—to make it clear that this is a context-dependent judgment. Furthermore, please position these examples not as recommendations to withdraw or reject the use of AI, but as case studies illustrating the decision-making process for appropriate review, suspension, and redesign within AI lifecycle management.
6	Regarding Sound Practice 9 (Performance management), the case study on page 39 describes a situation where, after comparing the performance of an AI model and a non-AI model, the organization decided to continue using the non-AI model. This suggests that AI is not one-size-fits-all and that decisions should be made by weighing factors such as each company's risk appetite. Conversely, examples of organizations that have accepted the risks associated with AI (such as accountability, transparency, hallucination, and bias) and continued to use it could serve as a useful reference for future implementations. Regarding Sound Practice 10 (Human oversight), the statement on page 44 reads: "As the use of AI agents across the financial institution increases, adapting human resource controls and processes to AI agents in a way that treats them as synthetic employees", rather than viewing AI agents uniformly as 'mere tools', they can be considered 'entities performing duties similar to employees' to a certain extent, particularly for the purpose of clarifying their operational roles, authorities, and lines of responsibility. However, simply using the metaphor of "synthetic employees" raises concerns that it might lead to the misunderstanding that AI agents must be equated with legally liable entities or employees under labor law. We would like to clarify that this statement does not imply that AI agents must be treated as entities subject to legal liability or as employees under an employment contract; rather, its intent is merely to suggest that, from the perspective of human resource management, it is conceivable to virtually include AI agents as resources.

	Since the examples provided focus primarily on the practices of large financial institutions, there is concern that they may be treated as the standard for supervisory expectations. As this could impose an excessive governance burden on AI used for low-risk applications or on small and medium-sized subsidiaries, we request that it be made clear that these examples are provided for illustrative purposes only.
7	On page 44, treating AI agents as "synthetic employees" is cited as an example of "human oversight". This can be interpreted to mean that synthetic employees could potentially replace humans. Regarding this point, we believe it would be more useful in practice if the definition of "AI Agent" or "agentic AI" in the Glossary not only classified AI agents as "AI systems that autonomously execute tasks", but also noted that, in the practical operations of financial institutions, they may be treated as "units of work execution similar to employees" in the design of controls such as authorization management, segregation of duties, supervision, and monitoring. In the Glossary, in addition to covering an AI Agent's autonomy, interaction with the external environment, use of tools, and impact on business processes, it would be desirable to supplement the definition by noting that, in the context of control design, it may be treated as a "business execution unit" subject to management—including identity and access management, authorization, approval workflows, activity logs, escalation procedures, and shutdown functions.