

FSB 市中協議文書「AI の責任ある導入のための適切な実践」 に意見提出

日本損害保険協会(会長：石川 耕治)は、金融安定理事会(FSB)が2026年6月10日から2026年7月22日にかけて市中協議に付した「AI の責任ある導入のための適切な実践」市中協議に対する意見を提出しました(詳細は添付1をご参照ください)。

<市中協議の概要>

- ・ 責任あるAI の導入により、金融機関は機会とメリットを最大限に活用しつつ、リスクを最小限に抑えることができる。特に、金融機関はAI がもたらす機会とリスクの理解、最新情報の把握、および適切な導入戦略とガードレールにより、変化し続ける関連リスクに対応する必要がある。
- ・ 本報告書は、金融システムにおけるAI の利用に伴うメリットとリスクを明らかにしており、金融機関が組織全体のAI ガバナンスおよびAI の開発・導入の各段階(AI ライフサイクル)の管理において適用できる12の健全な実践手法(sound practice)を提案している。また、金融機関による実際のAI 導入事例に基づくケーススタディが含まれている。
- ・ これらの指針は、AI がますます普及する環境下において、金融機関の取締役会や経営陣が事業戦略、技術導入、リスク管理を検討する際の指針となることを目的としている。

当協会は、FSB における国際保険監督基準策定の議論に積極的に参加しており、今後も市中協議等に際して本邦業界の意見を表明していきます。

(ご参考) 市中協議文書リンク：

<https://www.fsb.org/2026/06/sound-practices-for-responsible-adoption-of-artificial-intelligence-ai-consultation-report/>

No.	質問(仮訳)	損保協会意見(和文)	損保協会意見(英文)
1	本報告書に記載されている、金融機関による AI 導入のメリットおよびリスクについて同意するか。本報告書で取り上げられていない実質的なメリットやリスクはあるか。	・報告書の内容について概ね同意する。そのうえで、以下を指摘する。 ・「はじめに (p.3)」で言及されている「AI を導入しないことのリスク」について、より明確な整理が有用だと考える。例えば、不正検知対策のように、AI を活用しないこと自体が監督上の期待に十分応えられないリスクとなり得る領域も存在する。「導入しないリスク」と「導入することに伴うリスク」をどのように比較衡量すべきかについての考え方が示されれば、金融機関が責任ある導入判断を行う上で一層参考になると考える。 ・また、AI の成果物が他社の知的財産権を侵害するリスクも考えられる。	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	これらの健全な実践は、金融機関による責任ある AI 導入を可能にするために、十分に包括的かつ明確であるか。	・Sound Practice1 (Strategic direction and oversight) に関して、p.19"The board and senior management consider AI risks when setting risk appetite and tolerance"の 1 段落目において、"automated decision-making in critical business areas"が禁止される AI のユースケースとして例示されている。これはあくまで例示であり、重要性・顧客影響・法規制・補完統制に応じた判断であることを明確化いただきたい。禁止される内容が明確ではなく、AI 利活用の観点から阻害要因となりうる。	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.

		・取締役会・経営陣が、AI 活用方針、リスクアペタイト、重要ユースケースの監督、エスカレーション体制等の枠組み設定に関与することは重要であり、原案の趣旨に賛同する。そのうえで、実務上の理解をより明確にする観点から、取締役会・経営陣の関与は、個別ユースケースの詳細判断を一律に求めるものではなく、重要性・リスク・顧客影響等に応じて、各業務部門および統制部門が適切に判断・管理できる余地を前提とするものであることが示されると有用と考える。	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	これらの健全な実践指針は、あらゆる形態の AI に関連するリスク管理と、GenAI やエージェント型 AI といった新興かつ複雑な形態の AI に関連するリスクへの対応との間で、適切なバランスを保っているか。	・ Sound Practices3 (Incorporation of AI risks into risk management framework)に関して、p.23 の「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.」との記載に関して、第三者 AI プロバイダー、特に基盤モデルやクラウド等を提供する大規模プロバイダーを利用する場合、金融機関として、可能な範囲で契約上の情報提供、変更通知、設定管理等を確保することは重要である。一方で、市場慣行等により、必要な情報や対応手段を十分に確保することが実務上容易でない場合も想定される。そのため、金融機関側で利用範囲の限定、出力検証、継続的モニタリング、問題発生時の報告・対応プロセス、必要に応じた利用停止・見直しの条件等を組み合わせてリスクを低減する考え方が示されると有	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

	用である。 ・ Sound Practice 10（Human oversight）に関し、本報告書はエージェント AI に対する human-in-command 等の監督形態を示す一方で、エージェントの利用がスケールするとリアルタイムの人的監視が非現実的になり得ることも認めている。この点を踏まえ、人による監督から AI による監督（AI-in-the-loop）へ移行については、AI による監視自体にもモデルリスク・説明可能性・責任分界の課題があるため、最終的な説明責任は人・組織が保持すること、監視 AI の独立性・性能検証・ログ記録・人へのエスカレーションを含めた実務例が示されると有用だと考える。加えて、human-in-the-loop のたがを置き換えていくことの社会的受容性に対する合意形成について、業界や監督機関による具体的に踏み込んだ整理が併せて示されることも、金融機関によるエージェント活用において非常に有用であると考ええる。 ・p.28 でリスクへの対応策として AI のデータへのアクセス権の制限について記載されているほか、p.54 で AI を利用できる人を限定すること（研修を受けた人など）などが挙げられているが、アクセス権の限定などは AI の効果を減少させることになりかねず、低リスクの生産性向上用途まで過度に制限しないよう、データ分類・用途別の段階的統制が望ましいと考える。	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 gradual replacement 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
--	--	---

		・ AI ユースケースの可視化やインベントリ管理は、AI リスクの特定・評価・モニタリング・管理に資する有用な手段の一つと考えられる。一方で、全ての AI 利用について同一水準の文書化やライフサイクル管理を求めることは、特に低リスクの生産性向上用途や限定的な内部利用において、実務上過度な負担となる可能性がある。そのため、AI インベントリや文書化の範囲・粒度については、重要性、リスク、顧客影響、第三者依存度等に応じて柔軟に設定できることが明確になると有用である。	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	これらの健全な実践は、新たな種類の AI や、時間の経過に伴う責任ある導入に対応し、対処するのに十分な柔軟性を持っているか。	・ 本報告書は sound practices のメニューと位置づけ柔軟性を担保するとされているものの、詳細な記載（特に SP10 の人的な監督など）は、監督当局によって事実上の遵守チェックリストとして運用される懸念がある。これが生じると意図された柔軟性の AI 活用の阻害要因となりうる。本 SP を「網羅的な遵守チェックリストとして運用しないこと」ことを明示していただきたい。 ・ p.43 の下から 2 ポツ目に関して、p.42 の"Human-in-the-loop"の説明を踏まえると、保険金請求への対応にも適用される可能性があると考えられる。AI の進化を踏まえると、保険の内容等によっては、今後、人的承認や二重承認が不要となるものもあるのではないかと考えられ、人的承認・二重承認を常に求めるのではなく、低リスク・少額・定型的な保険金支払等	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.

		では、事前承認に代えて事後モニタリング、サンプリング、閾値超過時の人手エスカレーション等を認める旨を明記してはどうか。GenAI・エージェント型 AI に関して記載されている具体的統制は、例示であり、実際には各社が重要性やリスクに応じた統制を採用できることを明確化していただきたい。	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	本報告書の事例研究は、異なる種類の金融機関が責任ある AI 導入からどのように利益を得られるかを十分に浮き彫りにしているか。報告書に追加すべき事例研究はあるか。ある場合は、特にノンバンクに関する事例研究を提供してほしい。	・事例研究は有用であるが、掲載されている事例について、定量的成果を伴う成功例に偏っており、AI の採用を見送った、あるいは導入を中止した判断に関する事例は限られている。責任ある導入の観点では、撤退・不採用の判断に至った事例についても共有があれば実践的な示唆に富むものと考えられる。 その際、撤退・不採用の事例紹介が、将来的な禁止や制約事項と誤認・解釈されるおそれもあるため、こうした特定の技術やユースケースの禁止を意図するものではないことが誤解なく伝わるよう、「当該ユースケースや技術類型が不適切である」という結論のみを示すのではなく、判断時点の前提条件、リスク評価の観点、代替的統制の有無、データ・モデル・業務プロセス上の制約、再検討可能性等を含め、文脈依存の判断であることを明確にしていきたい。 また、撤退・不採用を推奨する趣旨ではなく、AI ライフサイクル管理における適切な見直し・停止・再設計の判断プロセスを示す事例として位置付けていただきたい。	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	本報告書の事例研究は、金融機関が AI を責任を持って導入する上で、実践的な示唆を提供しているか。	・ Sound Practice9 (Performance management) に関して、p.39 の事例において、AI モデルと Non-AI モデルのパフォーマンスを比較した結果、non-AI モデルの使用を継続したとのケースが記載されており、AI が万能ではなく、個社のリスクアベタイトなどとの比較で判断すべきという示唆になっている。これとは逆に、AI のリスク（説明責任、透明性、ハルシネーション、バイアスなど）を受容して使用した例などがあれば、今後の導入の参考になるのではないか。 ・ Sound Practice 10 (Human oversight)に関して、p.44 の"as the use of AI agents across the financial institution increases, adapting human resources controls and processes to AI agents in a way that treats them as synthetic employees."との記載について、AI Agent を一律に「単なるツール」として捉えるのではなく、業務上の役割・権限・責任分界を整理する目的においては、一定程度「従業員に類似した業務遂行主体」として捉えることは有用だと考えられる。ただし、synthetic employees という比喻を用いるだけだと、AI Agent を法的責任主体や雇用上の従業員と同一視しなければならないかのような誤解を招く懸念がある。本記載は、AI Agent を法的責任主体や雇用上の従業員として扱わなければならないという趣旨ではなく、あくまで人的リソース管理の観点で AI エージェントを仮想的にリソースとして含めることも	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.

		考えられるという趣旨であることを確認したい。 ・記載の事例は大規模な金融機関の実務を中心としているため、記載の事例が監督上の期待水準として扱われることが懸念される。低リスク用途の AI や中小規模子会社に対して過度なガバナンス上の負担が及ぶおそれがあるため、事例はあくまで例示であることを明確化していただきたい。	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	用語集の定義は明確であり、AI の最近の動向を含む業界のベストプラクティスと整合しているか。	・p.44 で AI agent を synthetic employees として扱うことが Human oversight の例として挙げられており、synthetic employee が human を代替することも想定され得ると読める。この点に関して、Glossary における AI Agent または agentic AI の定義については、AI Agent を「自律的にタスクを実行する AI システム」として整理するだけでなく、金融機関の実務においては、権限管理・職務分掌・監督・モニタリング等の統制設計上、「従業員に類似した業務遂行単位」として扱われ得る点も補足すると、より実務的に有用と考える。Glossary においては、AI Agent の自律性・外部環境との相互作用・ツール利用・業務プロセスへの影響に加え、統制設計上は、ID・アクセス管理、権限付与、承認フロー、行動ログ、エスカレーション、停止機能等の管理対象となる「業務遂行単位」として扱われ得ることを補足することが望ましい。	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.