



FINANCIAL
STABILITY
BOARD

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

Response to Consultation

Wise

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?

Wise agrees broadly with the benefits and risks identified in the consultation. Our comments below focus on the final part of the question, "Are there any substantive benefits or risks not covered?"

The FSB report is comprehensive across the AI risk landscape. It covers governance gaps, data quality, explainability, cyber risks, third-party dependencies, and agentic AI. However, some issues specific to payments, and particularly to cross-border payments and non-bank payment service providers, would warrant deeper treatment. We highlight three areas below.

A. Machine-to-machine payments and the shift from KYC to Know Your Agent ("KYA")

The consultation discusses agentic AI as a source of risk within financial institutions. It says less about what happens when AI agents become the customers or act on behalf of customers in payment flows.

We are moving toward a reality in which AI agents will compare prices, choose products and services, and initiate payments. When an AI agent transacts on behalf of a human or a firm, the counterparty is an autonomous digital actor that can sustain multiple interactions without constant human or institutional guidance or supervision.

In fact, the foundations of this machine-to-machine economy are already being built through the emergence of open protocols like Agent-to-Agent (A2A) negotiation and the Model Context Protocol (MCP). As these standards mature, the transition to new compliance frameworks must accelerate to keep pace.

This scenario creates a set of questions the FSB report does not fully address. If a payment is initiated by an AI agent, how do we verify that the agent is legitimate, properly authorised by its principal, and acting within scope? Traditional KYC frameworks usually assume a natural person or legal entity at the other end of the transaction that is performing sequential actions. These frameworks might not fully apply to autonomous digital actors that lack formal documentation about who they are and that may be handling hundreds of simultaneous

requests. In the face of these challenges, identity verification and session authentication need to become continuous, millisecond-fast, and based on reusable digital credentials.

These questions can have direct implications for the proposed sound practices. Sound Practice 5 (materiality and risk assessment) and Sound Practice 10 (human oversight) focus on the internal perspective of agentic AI. The proposed risk framing captures agents deployed within a financial institution's environment. It does not fully capture how financial institutions respond to interactions initiated by external agents, particularly those deployed by customers. A complementary view is needed for this situation.

B. Consumers "without eyes" and the need for machine-readable transparency

The report addresses consumer protection risks in section 1.3, including inadequate disclosure and unfair treatment. It does not directly address a related issue: that consumers could increasingly interact with financial services through AI intermediaries rather than directly.

As AI agents evolve to evaluate, recommend, and choose financial providers on behalf of consumers, disclosure designed for human eyes becomes insufficient. AI agents cannot compare the true cost of products or services if pricing information is not available in a structured, queryable, and machine-readable way.

And the risk here is not theoretical. According to industry reports, automated bot traffic surpassed human traffic on the internet for the first time in 2024 and continues to outpace it at an unprecedented rate. The web is moving from a place primarily inhabited by humans to one where AI agents make up a significant and growing share of the audience.

For financial services, this transition suggests that transparency and disclosure obligations must also evolve. Price and fee data should be published in structured formats that AI agents can reliably retrieve to evaluate costs correctly and make fair comparisons. Only then will it be possible to maintain a credible commitment to price transparency in payment services and to ensure that existing consumer protection rules remain effective even when consumers delegate decisions to machines.

C. Instant decisioning and the limits of human oversight

Sound Practice 10 (human oversight) sets out a useful taxonomy of human oversight types for AI actions. The report notes that continuous human monitoring becomes impractical as AI use scales.

In payments, this tension can be even more severe, especially in the face of instant payment systems that process transactions in seconds. When financial institutions allow AI tools to be deployed or used on instant rails, unexpected model behavior or malicious interventions can cause automated actions to quickly become unmanageable. This rapid execution compresses the timeline available for human intervention, creating challenges for traditional risk frameworks.

Similarly, payment fraud that moves at machine speed easily outpaces rules-based detection methods combined with human review. Wise's experience with real-time liveness checks, highlighted below, illustrates this problem. When an account is flagged for

suspicious activity, reducing the average resolution time for false positives from six hours to five minutes required removing human intervention from the initial verification loop. Requiring early human intervention would have preserved a form of oversight but with the increased likelihood of extending the suspension of legitimate customers and slowing responses to actual fraud.

This experience offers an insight that the proposed sound practices could reflect more clearly. In some payment use cases and in high-velocity fraud environments, the reduced opportunity for manual review workflows is not a governance failure but a necessary technical feature.

In any case, prioritizing AI intervention must follow clear standards, especially when it is introduced into customer-facing journeys. The AI system must meet quality thresholds before deployment and its decisions need to remain auditable. Moreover, meaningful human oversight must be preserved at the design, escalation, and post-deployment monitoring stages. At Wise, before any AI feature is rolled out, it is thoroughly tested to ensure that it is at least as accurate and effective at resolving issues as the current solutions in place. If it does not meet these standards, we reevaluate its deployment.

The FSB could further emphasize that the opportunity for human oversight will depend on the speed at which it must be taken. Making this relationship more explicit would help authorities and institutions better understand the challenges involved in reaching effective decisions in real-time environments.

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?**

Based on the issues highlighted on Wise's response to Question 1, we suggest that the FSB consider whether an additional sound practice, or an expansion of Sound Practices 5, 10, and 11, would help reflect specific issues related to machine-to-machine transactions. The framing could be:

"Sound practice 13 (Machine-led financial transactions): Financial institutions prepare for scenarios in which AI agents become dominant in financial transactions, including by verifying agent legitimacy and authorisation, maintaining continuous identity trust scoring, and defining agent-specific access, session, and audit controls."

Wise recognises this area is still developing. But adding a separate sound practice now, even in an exploratory form, would signal to the financial sector that standard-setting bodies and financial authorities are alert to the transformation and would encourage institutions to build the necessary controls before this reality becomes more widespread.

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.

We offer the following case studies from Wise to illustrate how non-banks can benefit from the responsible adoption of AI. Each case is tied to one or more of the sound practices proposed in the consultation.

Case study 1. Real-time liveness checks for false positive resolution

Relevant sound practices: Sound Practice 9 (performance management) and Sound Practice 10 (human oversight).

When a Wise account is flagged for suspicious activity, the customer can opt in to authenticate themselves through the Wise app in under 5 minutes on average. In the process, the customer verifies some personal data and takes a selfie. AI-powered facial recognition and liveness detection check the customer against their existing profile, unless the system recognizes one of the circumstances that would require human intervention, as explored in Case study 4 below.

Before deploying this functionality, resolving the same issue through customer support, manual checks, and several interactions typically took 6 hours from the customer's initial contact. The AI-enabled process reduces resolution time by 98.6%. This real-time liveness workflow serves as Wise's foundational step in line with a broader industry trend: making biometric authentication and 'zero-UI' session verification, which replaces traditional passwords with continuous contextual authentication, more common for everyday transactions.

This case study confirms Sound Practice 9. Performance thresholds were set before deployment, ongoing monitoring is in place, and the model is measured against a clear accuracy benchmark. It also refines the reading of Sound Practice 10. Removing the human from the initial verification loop improved consumer outcomes, avoiding prolonged suspension of legitimate customers. Meaningful human oversight is preserved at the design and post-deployment stages, but not in every individual case.

The wider point here is that in fraud prevention, faster does not mean weaker. Resolving false positives more quickly and effectively improves detection capabilities and frees operational capacity to prevent actual fraud.

Case Study 2. AML Payment Pattern Analysis and Proof-of-Address automation

Relevant sound practices: Sound Practice 6 (selection) and Sound Practice 9 (performance management).

In transaction monitoring, Wise's AML Payment Pattern Analysis has reduced handling time by 12%. In proof-of-address checks, AI-driven automation has achieved 95% automation at 99.4% accuracy. These are not incidental gains. In payments, speed and accuracy at the point of onboarding and monitoring are crucial. Faster decisions with lower friction allow

institutions to grow their customer base more consistently, while improving decision accuracy helps maintain, if not improve, compliance quality.

This case study confirms Sound Practice 6 on selection. Both deployments use vertical, workflow-embedded models chosen for a specific regulated task, rather than generic foundation models. For high-materiality compliance tasks, a purpose-built model with clear performance benchmarks might outperform a general-purpose alternative.

It also confirms Sound Practice 9 on performance management. Both use cases operate against explicit performance thresholds and are monitored on an ongoing basis. The 99.4% accuracy figure for proof-of-address illustrates the point made in Sound Practice 9 that institutions should specify thresholds ante hoc and review performance against them.

Case study 3. LLM Gateway with centralised logging, three lines of defence, and Consumer Duty Impact Assessments ("CDIAs")

Relevant sound practices: Sound Practice 2 (governance and accountability), Sound Practice 3 (incorporation of AI risks into risk management framework), Sound Practice 11 (cyber and ICT risk management), and Sound Practice 12 (third-party AI risk management).

Wise routes its production AI traffic through a centralised LLM Gateway. Inputs and outputs are logged with privacy, security, and first-line and second-line-of-defence approvals built in. Wise operates a three-lines-of-defence model, with Model Risk Management, CDIAs for customer-facing features, and alignment with the existing principles and rules from different jurisdictions.

Wise is deliberately model-agnostic across the frontier providers. This approach mitigates concentration risk at the model layer. It also allows us to route different workloads to the most suitable model based on cost and accuracy.

This case study confirms Sound Practices 2, 3, 11, and 12. It offers a concrete non-bank example of how a centralised control gateway, combined with a multi-vendor model strategy, addresses the concentration risk the FSB flags in Section 1.3.

Case study 4. Human escape hatches for vulnerable customers

Relevant sound practices: Sound Practice 10 (human oversight) and the consumer protection considerations in Section 1.3.

Wise operates a hard "Extra Care required" routing that sends vulnerable customers directly to human agents, bypassing AI-driven flows. The trigger is applied through both explicit customer signals and inferred vulnerability markers from customer interactions. This approach reflects two standards that Wise applies across customer-facing AI: transparency about when AI is involved and what actions it can take, and human intervention when the customer has to make material decisions or is in a vulnerable position.

This case study confirms Sound Practice 10 and reinforces the point made under Question 1. The correct form of human oversight depends on the customer, the use case, and the materiality of the interaction. Institutions need the flexibility to design oversight around the specific risks they face and the customers they serve.

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?**

The glossary is a useful foundation and captures the core AI concepts. Our comments below are limited to gaps and refinements that follow directly from the arguments and case studies from Wise set out under Questions 1, 3, and 5.

Suggested additions

"Agent authentication" and "Know Your Agent" ("KYA"): Both terms are central to the argument under Question 1 that traditional Know Your Customer ("KYC") frameworks cannot be fully applied to autonomous agents transacting on behalf of customers or firms. A working definition could describe "agent authentication" as the process by which a financial institution verifies the identity, legitimacy, and scope of authority of an AI agent initiating an interaction or transaction. "KYA" could, in turn, describe the broader set of controls financial institutions apply to establish and maintain trust in external AI agents, including verifying their principal, credentials, and session integrity.

"Consumer" or "End-user": Given the analysis of consumer protection risks in Section 1.3 of the consultation and our argument under Question 1 that consumers might increasingly interact with financial services through AI intermediaries, a working definition would be useful. It could describe "consumer" or "end-user" as the natural or legal person receiving a financial service, whether interacting directly with the financial institution or through an AI agent acting on their behalf.

Suggested amendments

"Agentic AI": The current definition mainly describes agents operating within a financial institution's environment. It does not capture agents acting outside institutional boundaries or on behalf of customers, which is a growing category of agentic AI explored in our response to Question 1. We suggest expanding the definition to reflect that agentic AI systems may operate within a single institution, across institutions, or on behalf of external parties, and that the risks differ across these scenarios.

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

Wise appreciates the opportunity to contribute to this timely discussion. We commend the FSB for producing a report that is technology-neutral, proportionate, and grounded in real practice. The 12 sound practices offer a menu that financial institutions of different sizes and levels of AI maturity can use to guide responsible adoption.

We particularly welcome the FSB's emphasis on public-private collaboration as a foundation for effective AI exploration. The consultation itself notes that public-private collaboration platforms, such as industry consortia and regulatory sandboxes, can build shared understanding of AI opportunities and risks. Wise shares this view. AI evolves faster than any single institution or authority can track on its own. Sharing information and knowledge

between regulators and regulated institutions is critical, both to keep pace with technology and to prevent ill-designed rules that might push innovation outside the regulatory perimeter.

As AI in financial services evolves, Wise remains committed to advancing effective practices that support responsible AI adoption while preserving customer protection and strengthening the global financial system.