



Jersey Financial
Services Commission

Guidance on the use of AI in Jersey's financial services sector

Issued: July 2026

Contents

1	Executive summary	3
2	Five principles for responsible AI use	4
2.1	Principle 1: governance and accountability.....	4
2.2	Principle 2: cybersecurity, privacy and data quality.....	4
2.3	Principle 3: regulatory compliance	6
2.4	Principle 4: consumer protection and fair treatment	7
2.5	Principle 5: transparency and explainability.....	7
3	Illustrative AI scenarios and mitigations	8
3.1	Scenario 1: AI supports AML onboarding with effective controls	8
3.2	Scenario 2: AI-assisted regulatory report with verified data	9
3.3	Scenario 3: third-party AI service outage.....	9
3.4	Scenario 4: AI fraud detection and model drift.....	10
3.5	Scenario 5: biased credit lending model.....	10
4	Glossary.....	11
	Appendix 1: Materiality ladder for AI use cases	13

1 Executive summary

This guidance explains how firms can use Artificial Intelligence (AI) responsibly in Jersey's financial services sector. It does not create new requirements. It aims to help firms apply their existing legal and regulatory obligations in a proportionate way.

We encourage firms to explore AI. Used responsibly, AI can improve efficiency, automate routine tasks and increase productivity. But AI also brings risks. Firms need to understand those risks and put suitable controls in place.

We have designed this to be risk-based and proportionate. For example, low impact productivity tools only require light touch controls unless they affect regulated activity, customer outcomes, confidential data or regulatory submissions. This guidance will be more necessary for higher impact use cases such as customer onboarding, lending decisions or investment advice.

The main principle is simple. Firms remain accountable for outcomes when they use AI, just as they do when they use any other technology. Where appropriate, firms should include material AI use cases and controls in their business risk assessment.

Our guidance closely aligns with the OECD AI Principles and guidance issued by international peer regulators.

This guidance is based on five principles:

1. **Governance and accountability:** Establish clear ownership, oversight and accountability for AI-assisted processes and outcomes.
2. **Cybersecurity, privacy and data quality:** Ensure that firms integrate AI into their existing cybersecurity arrangements whilst maintaining data quality and compliance with their data protection obligations.
3. **Regulatory compliance:** Ensure AI use is consistent with regulatory obligations and does not weaken compliance outcomes.
4. **Consumer protection and fair treatment:** Ensure AI supports fair customer outcomes and does not introduce bias or harm.
5. **Transparency and explainability:** Be able to explain if asked, to a reasonable degree commensurate with risk and complexity, AI-assisted decisions and outcomes.

2 Five principles for responsible AI use

These principles are linked and should be applied with regard to the risk and impact of the AI use case.

2.1 Principle 1: governance and accountability

As with other forms of technology, the task for boards and senior managers is to oversee integration into the business in a way that maintains accountability for decisions, conduct and outcomes. For AI, this may include human oversight of higher-impact decisions, challenge of model outputs and board reporting on material use cases. The firm remains accountable for compliance with all relevant laws and regulatory requirements when it uses AI.

We have included a guide to materiality at Appendix 1.

For material AI use cases, firms should consider proportionate measures such as:

- › setting clear governance arrangements for approving and reviewing material AI use cases, including approval thresholds and escalation routes
- › integrating AI into existing governance and control frameworks, in a way that is proportionate to risk
- › maintaining an AI register for each material AI use case
- › testing AI systems before and after deployment
- › applying stronger controls where AI systems can take actions, initiate workflows, access tools or carry out recommendations without real-time human approval ('agentic AI')
- › maintaining appropriate human oversight, especially for compliance-critical or customer-impacting decisions
- › ensuring that they have sufficient people with the required skills and capabilities to be able to apply human oversight and deliver accountability

2.2 Principle 2: cybersecurity, privacy and data quality

Firms with material AI use cases would benefit from robust operational controls, including cybersecurity, data protection, data quality, third-party oversight, and effective record-keeping.

2.2.1 Cybersecurity

AI systems should be integrated into the firm's existing cybersecurity processes rather than treated as a separate concern. For example:

- › use strong authentication, including multi-factor authentication where appropriate, particularly where AI systems access sensitive data or can take actions that affect customers, transactions or compliance
- › encrypt data at rest and in transit, including data sent to and received from third-party AI services
- › include AI systems in cybersecurity monitoring and incident response planning. Consider what AI-related incidents may require escalation under existing obligations, such as data leakage, incorrect regulatory submissions or missed AML red flags

For higher-impact or more complex AI deployments, firms may also wish to consider:

- › AI-specific threats, such as prompt injection, data poisoning and attempts to extract confidential information from AI outputs
- › segregation of development, test and production environments
- › content controls for customer-facing tools (for example, controls to reduce the risk of misleading, inappropriate or confidential outputs)

2.2.2 Data protection and privacy

Firms should ensure AI use complies with the Data Protection (Jersey) Law 2018. Key considerations include:

- › establishing and recording a lawful basis for processing personal data through AI systems
- › not inputting sensitive data into third-party tools unless approved and appropriately protected
- › conducting data protection impact assessments for use cases that handle personally identifiable information
- › applying data minimisation and appropriate safeguards
- › managing cross-border data transfers lawfully
- › updating privacy notices where AI materially changes how personal data are processed

2.2.3 Data quality

AI systems are only as reliable as the data they use:

- › review and validate data used as input into AI models for accuracy and relevance, appropriate to the use-case. This should also reduce the risk of skewed outputs or unintentional bias
- › for AML and customer due diligence use cases, ensure that sanctions lists, watchlists, and other reference data are current, and validate source documents

2.2.4 Third party and outsourcing

Our Outsourcing Policy applies to outsourced activity that uses AI. We are not creating additional requirements; the following considerations apply existing guidance where firms use third-party AI services, proportionate to the materiality of the use case:

- › any material dependency created by the AI service
- › supplier due diligence on security, data protection, testing and resilience
- › contractual arrangements for oversight, access to information and incident notification
- › use of sub-contractors, or cross-border transfers of data
- › fallback or exit plans if the service is unavailable, unreliable or withdrawn

2.2.5 Logging and auditability

Where AI is used in regulated processes, firms should keep sufficient records to support auditability and effective challenge.

2.3 Principle 3: regulatory compliance

Firms remain responsible for AI used in regulated activities. They should have appropriate standards, oversight and controls, and should be able to check that outcomes are accurate, fair and open to review.

This is particularly important where firms use AI in compliance functions, customer communications or regulatory reporting.

2.3.1 Using AI in compliance processes

Where AI assists with compliance functions (such as customer due diligence, transaction monitoring, or sanctions screening), firms should ensure that they:

- › **maintain or improve standards:** calibrate AI thresholds and criteria to an appropriate standard
- › **apply human oversight:** use human review proportionate to risk. A person should remain responsible for final decisions in higher-risk cases and for identifying anomalies
- › **monitor performance:** monitor false negatives and false positives, and recalibrate where performance degrades (for example, due to changing fraud typologies)

2.3.2 Using AI in customer-facing roles

Where AI generates or assists with customer-facing communications (for example chatbots, automated advisers, investor reporting or drafting tools), we recommend that firms:

- › make clear when customers are interacting directly with an AI system, or where AI influences a decision affecting them
- › ensure outputs are accurate, fair and not misleading, regardless of whether a person or an AI generated them
- › apply review processes proportionate to the significance, scale and risk of the communications

2.3.3 AI in regulatory reporting and interactions

The requirement to ensure the accuracy of regulatory submissions, regardless of how they are produced, remains. We recognise that AI can help compile data, spot patterns and draft text. Keep quality assurance proportionate to the submission's importance. Consider training staff to verify outputs against source data and to spot fabrication.

2.4 Principle 4: consumer protection and fair treatment

We recommend that firms consider how best to ensure that AI supports fair customer outcomes and does not result in harm or disadvantage to individuals or groups.

2.4.1 Fair outcomes

Firms should:

- › identify and mitigate bias introduced through training data or model design (for example, models trained primarily on data from certain jurisdictions may produce skewed assessments)
- › test AI outputs for fairness in decisions that affect customers, including periodic reviews for consistency across comparable groups

2.4.2 Protecting customers

Firms should ensure AI:

- › does not exploit customers' lack of understanding or vulnerability
- › does not create opaque decision-making that customers cannot understand or challenge
- › supports, rather than undermines, good customer outcomes

2.4.3 Challenge, complaints and recourse

Firms should ensure customers affected by AI-assisted decisions can access the firm's usual complaints processes. This includes:

- › explaining how customers can request human review where appropriate
- › ensuring complaints handling procedures can deal with AI-related queries and challenges
- › training staff to handle AI-related complaints and to explain, in plain terms, how they reached an AI-assisted decision

2.5 Principle 5: transparency and explainability

Firms do not need to understand every technical detail of an AI model. However, they should be able to provide a clear description of the main factors that drove the outcome.

This should be supported by an up-to-date register of AI use cases, a risk-based assessment of their potential impact, and governance and controls proportionate to risk.

2.5.1 Customer-facing transparency

Where a customer interacts with an AI system (for example, a chatbot), firms should disclose this clearly.

For significant decisions, firms should consider what additional information customers need, including rights to further information, human review or complaint.

2.5.2 AI register and risk classification

Firms should keep an AI register for use cases that materially affect regulated activity, customer or client outcomes, financial crime controls, or their risk profile. The register should record:

- › what the AI system is used for
- › the type of AI used, in plain terms
- › the main data inputs
- › who is responsible for oversight
- › a risk/impact classification

When classifying AI use cases, firms should consider:

- › whether the AI directly affects customers or financial transactions, or is purely an internal tool
- › whether an output could lead to denial of a service, closure of an account, a blocked transaction, or a regulatory submission
- › the impact if the AI produces an incorrect output

Firms should use this classification to inform the firm's approach to governance, testing, monitoring and record-keeping. Higher-impact use cases warrant stronger controls.

3 Illustrative AI scenarios and mitigations

The following simplified scenarios illustrate how AI can be used responsibly.

3.1 Scenario 1: AI supports AML onboarding with effective controls

Case study

A firm uses an AI tool to review documents during customer onboarding. The tool helps identify missing information, unusual documents and possible red flags.

Why this may be appropriate

The AI supports the onboarding process. It does not replace the firm's AML obligations or the need for human judgement. The firm remains responsible for the final onboarding decision.

Requirements and assumptions

This use may be appropriate where the firm:

- › approves the use case through its governance process
- › records the tool in its AI register
- › tests the tool against known red-flag cases
- › keeps sanctions lists, watchlists and reference data up to date
- › routes higher risk cases to a suitably skilled person
- › records key inputs, outputs and human overrides
- › has fallback arrangements if a third-party service fails
- › ensures a competent person remains responsible for high impact or unusual cases

3.2 Scenario 2: AI-assisted regulatory report with verified data

Case study

A firm uses generative AI to help organise and draft the narrative sections of a regulatory report. The AI does not create figures or unsupported factual claims. Figures and factual statements come from approved internal source data.

Why this may be appropriate

The AI is being used to improve structure, language and efficiency. It is not being used as an unchecked source of facts. The firm remains responsible for the accuracy and completeness of the submission.

Requirements and assumptions

This use may be appropriate where the firm:

- › uses an approved AI tool or environment
- › restricts sensitive data from unapproved third-party tools
- › requires staff to verify AI outputs against source data
- › trains staff to identify fabricated or unsupported content
- › records reviewer sign-off
- › applies quality assurance that reflects the importance of the submission

3.3 Scenario 3: third-party AI service outage

Case study

A firm uses a third-party AI service to help interpret onboarding documents. The service becomes unavailable. As a result, potential red flags are not surfaced, but onboarding continues as normal.

Issue

The firm may onboard customers without appropriate checks. This may weaken AML controls and create operational resilience concerns.

Controls firms may consider

- › treat critical AI services as important outsourcing or third-party dependencies
- › agree suitable service levels and incident notification requirements
- › monitor service availability
- › pause, limit or manually review onboarding when the tool is unavailable
- › maintain manual fallback procedures
- › review whether the outage needs escalating under existing incident processes

3.4 Scenario 4: AI fraud detection and model drift

Case study

A firm uses an AI model to help detect fraud. At first, the model performs well. Over time, fraud patterns change. The model becomes less effective, but the firm does not identify the decline until after a significant fraud incident.

Issue

The firm may rely on a control that no longer works as expected. This could lead to missed fraud, customer harm or weakened financial crime controls.

Controls firms may consider

- › monitor performance using agreed metrics
- › track false positives and false negatives
- › set thresholds for review or escalation
- › periodically recalibrate or retrain the model
- › keep version control records
- › increase human review when performance drops
- › use rules-based or manual fallback controls if needed

3.5 Scenario 5: biased credit lending model

Case study

A firm uses an AI model to support lending decisions. Review of the model suggests that comparable applicants may receive different outcomes for reasons that are not explained by objective credit factors.

Issue

The model may produce unfair outcomes. This could happen because of biased training data, poor feature selection, weak governance or feedback loops over time.

Controls firms may consider

- › test outcomes for fairness and consistency
- › review whether the model uses inappropriate or unreliable features
- › assess whether training data are representative and relevant
- › ensure customers can challenge decisions

4 Glossary

Artificial intelligence (AI)	For this guidance, AI means a machine-based system that uses inputs to generate outputs. These outputs may include predictions, content, recommendations or decisions that can affect digital or real-world environments. This guidance is most relevant where AI materially affects regulated activity, a regulatory obligation, or a decision affecting clients or counterparties. Established statistical models already covered by existing regimes, such as solvency II internal models, IFRS 9 expected credit loss models and routine actuarial practice, are outside this definition.
AI register	A central record of the firm's material AI use cases. It may record what the AI is used for, who is accountable, the data it uses, the system or model involved, its risk rating, human oversight arrangements, any third-party supplier, and when it was last reviewed.
AML (Anti-money laundering)	Laws, regulations, and procedures that aim to stop criminals disguising the proceeds of crime as legitimate funds, and to detect this where it happens.
Auditability	The ability to show, after the event, how an AI-assisted outcome was produced. This may include the inputs used, the system version, the output generated, and any human review or override. Auditability is about reconstructing what happened. Explainability is about understanding why the system reached the result.
Bias (algorithmic)	A pattern in an AI system's outputs that unfairly favours or disadvantages particular people or groups. Bias can come from the training data, the model design, the way the system is used, or feedback loops over time.
Data minimisation	The principle that only the minimum personal data necessary for a specific, lawful purpose should be collected and processed, and held no longer than needed. A requirement under the Data Protection (Jersey) Law 2018.
Data poisoning	An attack where someone deliberately corrupts the data used to train an AI system. This can cause the system to behave incorrectly or produce unreliable outputs. It is a particular risk where training data comes from external or unverified sources.
Explainability	The ability to explain how an AI system reached a result, in a way that is suitable for the audience. This includes explaining the main

	inputs or factors that influenced the result. Explainability supports review, challenge and accountability.
Hallucination (in AI)	A fabricated or unsupported output produced by a generative AI system. This may include invented facts, figures, or citations, often presented in a fluent and confident way.
Human oversight	Review, challenge, approval, or intervention by a competent person in respect of an AI-assisted process or decision. The level of oversight should match the risk and impact of the use case. It may range from periodic checks of low-risk outputs to mandatory human approval before action is taken.
Material AI use case	An AI use case where failure, misuse, or poor performance could have a significant effect on customers, regulated activity, regulatory reporting, financial crime controls, or the firm's risk profile. Materiality should be assessed in context, including the scale and value of decisions, whether outcomes can be reversed and the strength of human checks.
Model	In AI, a model is the trained system that takes inputs and produces outputs. How it behaves depends on both its design and the data used to train it.
Model drift	A decline in an AI model's performance over time. This can happen when the real-world data it processes becomes different to the data it was trained on. Firms can manage this through monitoring, re-validation and re-calibration or retraining where needed.
Prompt injection	An attack on a generative AI system using crafted instructions. These can be entered by a user or hidden in documents, web pages or other content processed by the system. They may cause the system to produce unintended outputs, bypass controls or reveal information it should not disclose.
Training data	The data used to build and calibrate an AI model. Its quality, representativeness and lawful origin directly affect the reliability, fairness and defensibility of the model's outputs.

Appendix 1: Materiality ladder for AI use cases

We have designed this ladder to help you apply proportionate governance to AI use cases. It is not a checklist or new taxonomy. Not every factor is relevant for every use case.

Level	When this level is likely to apply	Factors to consider	Possible governance response
Low impact	The AI use case is narrow and unlikely to create material customer, regulatory or control consequences if it is wrong	› Customer impact: none or indirect only › Regulatory reporting: no impact › Financial crime: no impact or peripheral only › Data sensitivity: public, non-confidential or low-sensitivity internal data › Autonomy: none › Complexity: low › Reversibility: easy to correct or stop › Third-party reliance: limited and non-critical	› Record in the AI register › Assign an owner › Carry out basic supplier and data checks › Provide user guidance › Test before use › Review periodically
Medium impact	The AI informs decisions or forms part of an important workflow, but does not determine material outcomes without meaningful human review.	› Customer impact: indirect but real, or direct with meaningful human review › Regulatory reporting: may support analysis, drafting or preparation, but not final reporting without review › Financial crime: supports detection, triage or monitoring › Data sensitivity: confidential data or some personal data › Autonomy: decision-support › Complexity: moderate › Reversibility: correction is possible, but may require effort › Third-party reliance: material vendor or outsourced model dependency	› Require formal internal approval › Document the risk assessment › Carry out structured testing or validation › Define human oversight › Document limits of use › Set an escalation route for incidents › Apply stronger third-party due diligence
Higher impact/material	The AI use case could materially affect customer outcomes, core operations, regulated activity, regulatory reporting or market-facing activity.	› Customer impact: direct, significant or hard to remediate. › Regulatory reporting: affects submissions, disclosures or prudential/compliance outputs. › Financial crime: used in AML/CFT, sanctions, fraud detection, transaction monitoring or investigation. › Data sensitivity: highly confidential, sensitive or large-scale personal data. › Autonomy: high, with limited real-time human intervention. › Complexity: high, opaque, adaptive or multi-component. › Reversibility: difficult to unwind quickly. › Third-party reliance: critical dependency on external provider, model or infrastructure.	› Require senior approval › Document the materiality assessment › Include independent validation or challenge › Apply stronger monitoring and change control › Maintain explicit fallback or manual override › Apply enhanced third-party controls › Reassess periodically › Be ready to evidence governance and records