

BOARD BRIEFING REPORT

EUROPEAN BANKING & FINANCIAL MARKETS

Where Is AI Allowed to Act Without Human Approval?

A Board Briefing Report on agentic AI, human oversight, decision rights and control boundaries in European banking.



AI is moving from recommendation to action.
The issue is no longer only model output.
The issue is the action boundary.

ASSUME THE FOLLOWING

An AI-supported workflow is no longer only producing text, summaries or recommendations. It can trigger a workflow step, update a record, escalate an alert, draft a client communication, prioritise a case, recommend a control response or prepare an action for execution.

This scenario is generic and hypothetical. It does not describe any specific institution.

THE QUESTION THAT FOLLOWS

Where is AI allowed to act?

Where must a human approve?

Where is the boundary documented, monitored and reversible?

Three Orientating Facts for the Board

01

Agentic AI changes the governance question

Traditional AI governance asks whether a model is accurate, explainable, monitored and controlled. Agentic AI adds another question: what is the AI allowed to do?

02

Capability is not permission

A model can be capable of acting autonomously. That does not decide where it is allowed to act. Permission must be determined by materiality, reversibility, customer impact, legal effect and risk appetite.

03

Human oversight must be designed

A strong answer is not “there is a human in the loop”. A strong answer shows where AI proposes, where AI executes only after approval, where bounded autonomy is allowed, and where autonomous action is not permitted.

Source: EU AI Act; GDPR Article 22; ESMA AI Statement; IOSCO AI Toolkit.

Control Boundary, Not Technology Capability

GREENDECISION-GRADE	AMBERPARTIALLY EVIDENCED	REDNOT DECISION-GRADE
AI action inventory exists and is current.	AI use cases are inventoried, but action rights are not fully mapped.	<i>“It only recommends.”</i>
Each AI-supported workflow has a defined autonomy level.	Human review exists, but is inconsistently evidenced.	<i>“A human is somewhere in the loop.”</i>
Human approval points are mapped by role and authority.	Technical permissions exist, but are not linked to business decision rights.	<i>“The workflow is internal.”</i>
Tool access, API access and workflow permissions are controlled.	Exception logs exist, but are not reviewed at governance level.	<i>“The user can override it.”</i>
Actions are logged with actor, authorisation, timestamp and outcome.	Business ownership exists, but not for autonomy boundaries.	<i>“The model has no final decision right.”</i>
Exception, override and stop mechanisms are monitored.	Stop mechanisms exist in design, but have not been tested.	<i>“The vendor controls the permissions.”</i>
Board reporting distinguishes assistive, advisory and action-triggering AI.		<i>“We will define human approval during rollout.”</i>
		<i>“This is only a productivity tool.”</i>

GREEN is not a human-in-the-loop label. **GREEN is an evidenced control boundary.**

Technology may enable autonomy. The business owns the action impact. **The board needs to see the boundary.**

■ CIO / CTO	Architecture, workflow integration, tool access, API permissions, system boundaries and logging.
■ Business Owner	Decision impact, workflow purpose, materiality, reversibility and operational consequences.
■ CRO / Operational Risk	Risk classification, tolerance for autonomous action, residual risk and escalation thresholds.
■ Compliance / Legal / DPO	AI Act, GDPR, conduct, data protection, automated decisioning and regulatory interpretation.
■ CISO / Technology Risk	Identity, access, authorisation, cyber risk, misuse scenarios and technical stop mechanisms.
■ COO	Operational execution, control handovers, process redesign and service continuity impact.
■ Internal Audit	Independent review of action boundaries, logs, approvals, stop mechanisms and governance evidence.

Ten Artefacts That Make AI Action Governable

- 01

AI Action Inventory
Shows which AI-supported workflows only suggest, which trigger actions and which may act within defined boundaries.
- 02

Autonomy Level Classification
Shows whether the use case is assistive, advisory, approval-based, bounded autonomous or monitored autonomous.
- 03

Action Boundary Map
Shows what AI may do, what it may not do and what happens at the boundary.
- 04

Human Approval Matrix
Shows which action requires approval, which role approves it and whether the approval is meaningful.
- 05

Tool / API Permission Register
Shows which systems, datasets, workflow steps, communication channels or external tools AI can access.
- 06

Decision Rights Map
Shows who can approve, expand, restrict, pause or terminate AI action rights.
- 07

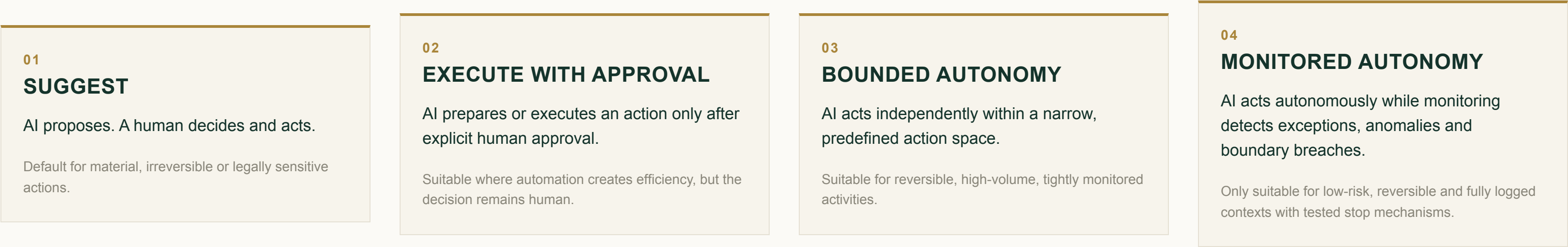
Exception and Override Log
Shows where AI actions were blocked, overridden, escalated or corrected.
- 08

Client / Conduct Impact View
Shows whether AI actions affect clients, advice, pricing, risk, service quality, complaints or regulatory outcomes.
- 09

Testing and Simulation Evidence
Shows what happens when AI misclassifies, escalates incorrectly, communicates incorrectly or acts outside expectation.
- 10

Stop / Rollback Protocol
Shows who can restrict, suspend or disable the AI capability — and whether that mechanism has been tested.

AI autonomy is not a binary switch. **It should be calibrated by materiality and reversibility.**



DESIGN RULE The autonomy level follows the materiality and reversibility of the action. It does not follow the capability of the model.

Warning Signals

These answers may sound reassuring in a board meeting. They are not yet decision-grade.

“It only recommends.”

A recommendation can still shape a regulated decision. The question is whether it changes behaviour, prioritisation or outcome.

“A human is in the loop.”

Human presence is not enough. The board needs to know whether the human has time, competence, authority and information to intervene meaningfully.

“The workflow is internal.”

Internal workflows can still affect risk assessments, controls, reporting, client outcomes or regulatory obligations.

“The user can override it.”

An override option is not the same as effective control. The board needs override frequency, review evidence and escalation logic.

None of these statements is necessarily wrong. **They are simply not sufficient evidence for board-level reliance.**

“The model has no final decision right.”

Final decision right is only one question. Triggering, ranking, escalating, suppressing or drafting can still materially shape the decision.

“The vendor controls the permissions.”

Technical permission management can sit with a vendor. Accountability for allowed action remains with the institution.

“We will define human approval during rollout.”

Approval boundaries are design requirements. They should not be discovered during deployment.

“This is only a productivity tool.”

A productivity tool can become a control issue when it touches regulated workflows, client information, risk views or decision rationales.

Six Conditions for an Evidenced Control Boundary

01 Action Rights Visible

The institution can show which AI-supported workflows can suggest, trigger, execute, communicate or modify records.

04 Permissions Controlled

Tool access, API access, workflow permissions and communication rights are approved, limited and monitored.

02 Autonomy Classified

Each AI use case is assigned an autonomy level based on materiality, reversibility and impact.

05 Actions Reconstructable

Every autonomous or semi-autonomous action can be reconstructed: what acted, under whose authority, on what input, with what result.

03 Human Approval Meaningful

Human approval points are designed so that the human can understand, challenge, override and stop the AI-supported action.

06 Stop Mechanism Tested

The institution can restrict, pause or terminate AI action rights quickly — and has tested that capability.

SUGGESTED NEXT QUESTION

Which AI-supported workflow would create the greatest accountability issue if it acted tomorrow without explicit human approval — and can we evidence the boundary that prevents that?

SELECTED SOURCES

EU AI Act, especially Articles 12, 14, 15, 19, 26 and Annex III.

GDPR, especially Article 22 on automated decision-making.

DORA, especially ICT governance, logging, detection and continuity requirements.

ESMA, Public Statement on AI and Investment Services, May 2024.

IOSCO, Supervisory Toolkit for AI Use in Capital Markets, 2026.

NIST AI Risk Management Framework and Generative AI Profile.

BOARD BRIEFINGS SERIES

Executive Questions for European Banking & Financial Markets

A human in the loop is not the same as human control.

A model can be capable of acting autonomously. That does not decide where it is allowed to.

www.monshizadeh.com
Reza Monshizadeh · Senior Technology Executive · Frankfurt