

Article 26 as Runtime Governance

Executive Summary / Public Brief

Evidence of Control, Human Oversight and Deployer Responsibility under the EU AI Act

Author	Jozsef Fodor
Project / Framework	NextOne Observatory
Version	V0.2 - Public Brief Draft
Date	19 June 2026
Document Type	Executive Summary / Public Brief

Core Operating Sentence

A deployer should be able to prove that a high-risk AI system was used under valid, monitored, human-controllable and reconstructable conditions when its output influenced consequence.

Core Thesis

Article 26 of the EU AI Act should not be treated only as a deployer compliance checklist. It can also be read as an operational governance problem.

For high-risk AI systems, the key question is not only whether a system was approved, documented, procured or registered. The harder question is whether the deployer can prove that the AI-supported workflow was still operating under valid control conditions when the system output began to influence real consequence.

Why This Matters

Many organizations already have AI policies, governance committees, risk assessments, system registers, training records and monitoring dashboards. These are important. But they do not automatically prove that an AI system was controlled when its output affected a person, workflow, process, decision or operational outcome.

A policy may say that human oversight exists. But was the human overseer trained, informed, available and authorized to intervene before the AI output became consequential?

A log may exist. But can it reconstruct why the AI-supported workflow was allowed to continue?

This is the gap between compliance language and operational proof.

From Static Compliance to Runtime Governance

Traditional compliance often focuses on documentation before deployment and audit after the fact. High-risk AI systems require something more dynamic. They require governance during use.

The conditions around an AI-supported workflow can change while the workflow itself still appears to function normally. Input data may become stale. A policy may change. A human overseer may become unavailable. A model may be substituted. A workflow may drift beyond its original intended purpose. An AI output that was supposed to be advisory may become practically decisive.

In those situations, the system may not visibly fail. The workflow may complete successfully. The dashboard may stay green. But the control conditions that justified continuation may no longer be valid.

Evidence of Control

Evidence of control is not just evidence that a policy exists. It is evidence that relevant control conditions were active, valid and inspectable when they mattered.

This may include evidence of system identity, intended purpose and actual use, provider instructions mapped to internal controls, human oversight competence and availability, input data relevance and freshness, monitoring status, alerts and escalation records, log retention, authority and approval state, incident response and reasons why the workflow was allowed, paused, escalated or refused.

The goal is not bureaucracy for its own sake. The goal is reconstructability. A mature deployer should not only know that an AI system produced an output. It should know why that output was allowed to matter.

Human Oversight Must Become Operational

The study argues that human-in-the-loop language is not enough. Human oversight is meaningful only when the human has a real ability to affect the outcome.

That means the human should have enough competence to understand the system's role and limitations, enough information to assess the output in context, enough time to review before consequence, enough authority to challenge or stop the workflow, and access to evidence, not only the final AI-generated output.

A human who can only approve what the workflow has already made inevitable is not truly in control.

The Four Runtime Governance States

ALLOW means the AI-supported workflow may continue because the relevant control conditions appear valid.

HOLD means continuation should pause because evidence, context, authority, data, oversight or scope is incomplete, stale or uncertain.

ESCALATE means the issue requires human, governance, risk, legal, technical, supervisory or senior review before continuation.

REFUSE means the system or workflow should not continue because the control conditions are invalid, prohibited, unsupported or unsafe for the intended consequence.

These states are not presented as legal categories. They are proposed as an operational governance model for turning deployer responsibility into practical control behaviour.

Final Conclusion

Article 26 is not only about AI compliance. It is about evidence-based runtime control.

The future of trustworthy AI will not depend only on better models, stronger policies or more complete documentation. It will also depend on whether organizations can prove that AI-supported workflows were allowed to continue only when the relevant control conditions were still valid.

The decisive question may not be: Did the AI system produce an output?

The decisive question may be: Why was that output allowed to matter?

Use Notice

This brief is provided for research, governance design and professional discussion purposes only. It is not legal advice. It does not replace legal, regulatory, compliance, technical, supervisory, audit, data protection or sector-specific advice.

Suggested Citation

Fodor, Jozsef. Article 26 as Runtime Governance: Evidence of Control, Human Oversight and Deployer Responsibility under the EU AI Act. NextOne Observatory, 2026.