

# Addressing the EU AI Act

---

## PlainID controls mapping for AI and agentic systems

How runtime authorization helps turn AI governance requirements into enforceable controls

**Discover • Manage • Authorize • Audit**

AI governance frameworks define what organizations should govern, document, monitor, and control. PlainID provides a runtime authorization control plane that helps operationalize those expectations at the moment humans, machines, and AI agents access data, invoke tools and APIs, or return sensitive information.

**Note:** PlainID can materially support implementation of governance and security requirements, but it does not by itself establish legal compliance, complete an AI risk program, or confer ISO certification.

## 1. Executive summary

The EU AI Act establishes a risk-based regulatory framework for AI systems and general-purpose AI. For organizations deploying enterprise AI and agentic systems, the practical challenge is not only documenting governance intent, it is ensuring that AI access and actions remain inside approved boundaries at runtime. PlainID helps translate policy into enforceable controls over identities, data, APIs, MCP tools and AI-generated outputs.

### Where PlainID is strongest

PlainID is particularly relevant to the EU AI Act requirements and deployer/provider obligations involving risk controls, access governance, record-keeping support, human oversight, cybersecurity, secure use, data exposure prevention and operational monitoring. It is complementary, not a replacement, for model testing, training-data governance, technical documentation, quality management, conformity assessment and fundamental-rights impact assessment.

### Current applicability context — August 2026

- The AI Act entered into force in August 2024 and is generally applicable from 2 August 2026, with staged exceptions.
- Prohibited-practice and AI-literacy provisions have applied since 2 February 2025; governance and GPAI obligations have applied since 2 August 2025.
- Following the 2026 simplification/Omnibus changes, many Annex III high-risk use-case rules apply from 2 December 2027; certain product-embedded high-risk systems apply from 2 August 2028.
- Organizations should still prepare control architecture now because inventory, policy design, integration and evidence collection require lead time.

## 2. PlainID-to-EU AI Act mapping

| Framework requirement / outcome                              | PlainID alignment                    | How PlainID helps                                                                                                                                                                                                                              | Coverage                                               |
|--------------------------------------------------------------|--------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------|
| <b>Article 9 — Risk management system</b>                    | Runtime risk treatment               | Enforces policy-based restrictions that reduce unauthorized access, excessive agency and sensitive-data exposure during operation; policies can incorporate identity, agent, resource and risk context.                                        | <b>Direct / strong</b>                                 |
| <b>Article 10 — Data and data governance</b>                 | Authorized data use                  | Restricts which data sources, rows, columns, cells or documents an AI workflow can retrieve based on the acting human/agent and context. Helps enforce approved use boundaries; does not validate training-data quality or representativeness. | <b>Supporting</b>                                      |
| <b>Article 11 — Technical documentation</b>                  | Control evidence                     | Policy definitions, enforcement architecture and decision evidence can form part of the technical control documentation package. PlainID does not create the full Annex IV technical documentation.                                            | <b>Supporting</b>                                      |
| <b>Article 12 — Record-keeping / logs</b>                    | Runtime decision traceability        | Authorization decisions and policy enforcement create evidence of requested access, applicable policy, decision and control action, complementing AI-system event logs.                                                                        | <b>Direct / strong</b>                                 |
| <b>Article 13 — Transparency to deployers</b>                | Explainable control boundaries       | Centralized, readable policy models clarify what an identity or agent is permitted to access or do and why access is restricted; this complements model/output transparency.                                                                   | <b>Supporting</b>                                      |
| <b>Article 14 — Human oversight</b>                          | Human-in-the-loop control            | Policies can require step-up, approval or a human checkpoint for high-risk actions and can prevent execution outside authorized scope.                                                                                                         | <b>Direct / strong</b>                                 |
| <b>Article 15 — Accuracy, robustness &amp; cybersecurity</b> | Cybersecurity / resilience of access | Least privilege, dynamic authorization, tool restrictions, parameter governance and output masking reduce attack paths and limit blast radius from compromised or misbehaving agents. PlainID does not establish model accuracy.               | <b>Direct for cybersecurity; adjacent for accuracy</b> |
| <b>Article 17 — Quality management system</b>                | Policy lifecycle governance          | Policy authoring, approval, promotion, versioning patterns and centralized management support repeatable operational controls within a broader QMS.                                                                                            | <b>Supporting</b>                                      |
| <b>Article 26 — Deployer obligations</b>                     | Controlled operation                 | Helps deployers use AI within approved conditions, assign scoped access, monitor authorization activity and retain relevant decision evidence.                                                                                                 | <b>Direct / supporting</b>                             |
| <b>Article 27 — Fundamental rights impact assessment</b>     | Mitigation implementation            | PlainID can implement access/data/action mitigations identified by an FRIA, but does not perform the rights-impact assessment itself.                                                                                                          | <b>Supporting</b>                                      |

| Framework requirement / outcome       | PlainID alignment     | How PlainID helps                                                                                                                                                     | Coverage   |
|---------------------------------------|-----------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------|
| Article 50 — Transparency obligations | Controlled disclosure | Output filtering/masking and policy-controlled interaction can reduce unauthorized disclosure; formal AI-content labeling and disclosure obligations remain separate. | Supporting |

### 3. EU AI Act control themes PlainID operationalizes

#### Know which AI agents exist

Discovery and classification reduce blind spots and enable a governed inventory of agents, connected tools and data sources.

#### Enforce least privilege continuously

Replace broad standing access with context-aware, task- and purpose-scoped authorization at execution time.

#### Constrain autonomous actions

Restrict tool/API invocation, MCP server and tool access, and relevant parameters before an action is executed.

#### Protect sensitive data

Apply row/column/cell filtering, document-level controls and output masking so AI receives and exposes only authorized information.

#### Preserve human control

Require approval or step-up for high-risk activities and deny execution when the required oversight condition is not met.

#### Create auditable evidence

Central policies and runtime decisions create a repeatable evidence trail for control testing, incident analysis and compliance review.

### 4. What PlainID does not replace

- AI system risk classification and legal applicability analysis under the Act.
- Training, validation and testing data governance, bias/representativeness assessment, or model accuracy validation.
- Full technical documentation, declarations of conformity, CE marking, registration or notified-body/conformity-assessment processes.

- Organization-wide AI literacy, workforce training, model cards, provider disclosures or GPAI model obligations.
- Fundamental-rights impact assessments or other legal/privacy impact assessments, although PlainID can implement resulting access-control mitigations.

## 5. PlainID capabilities used in the mappings

| Capability                                | Customer value                                                                                                                           |
|-------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------|
| <b>AI and agent discovery / inventory</b> | Visibility into AI agents, MCP tools, RAG/data sources and relationships; supports classification and governance onboarding.             |
| <b>Centralized policy management</b>      | Policy-based authorization with reusable, centrally governed policies across applications, APIs, data platforms and AI/agentic flows.    |
| <b>Composite identity and context</b>     | Runtime evaluation of human identity, agent/non-human identity, purpose or intent, resource sensitivity, environmental and risk signals. |
| <b>Input guardrails</b>                   | Restrict prompts or requested actions based on identity, agent, purpose and context before downstream execution.                         |
| <b>Data retrieval guardrails</b>          | Apply least-privilege restrictions, row/column/cell filters, and contextual data controls before data is retrieved or exposed.           |
| <b>MCP/tool/API guardrails</b>            | Control which services, MCP servers/tools and APIs may be invoked and constrain actions/parameters based on runtime context.             |
| <b>Output controls</b>                    | Mask or redact sensitive information and validate output against policy before delivery.                                                 |
| <b>Auditability and policy lifecycle</b>  | Central policy administration, approval/promotion patterns, observable decisions and traceable enforcement evidence.                     |

## 6. Reference implementation pattern

A common architecture pattern is to keep governance intent centralized while enforcing it at the points where AI systems actually act.

|                       |                                                      |                                                                                        |
|-----------------------|------------------------------------------------------|----------------------------------------------------------------------------------------|
| <b>1. Signals</b>     | SSO / directory / IGA / risk / agent metadata        | Identity, entitlement, assurance, device, agent, task and environmental context        |
| <b>2. Policy</b>      | PlainID centralized authorization policy             | Business rules translate governance intent into executable access decisions            |
| <b>3. Runtime</b>     | Input • Data • MCP/Tools • APIs • Output             | Decisions occur at the moment an AI/agent workflow requests access or action           |
| <b>4. Enforcement</b> | Authorizers / gateways / frameworks / data platforms | Allow, deny, filter, mask, constrain parameters, require step-up or human approval     |
| <b>5. Evidence</b>    | Decision logs / policy lifecycle / observability     | Traceability for monitoring, audit, review, incident investigation and control testing |

## 7. Recommended customer implementation approach

### 1. Scope and classify

Identify AI systems, agents, use cases, affected business processes and likely AI Act roles/risk categories.

### 2. Discover access paths

Inventory agent frameworks, MCP servers/tools, APIs, RAG sources, databases and downstream systems.

### 3. Define control policy

Translate legal/risk requirements into explicit rules covering who/what may access which resource, for what purpose, under which conditions.

### 4. Enforce at runtime

Insert authorization at prompt/input, retrieval, tool/API invocation and output points; remove unnecessary standing privileges.

### 5. Add oversight tiers

Define actions that require deny, step-up, human approval, masking, filtering or restricted parameter sets.

### 6. Validate evidence

Test policies, verify logs and review authorization outcomes as part of control assurance and post-deployment monitoring.

#### Key takeaway

The EU AI Act raises the bar from documented AI governance to demonstrable operational control. PlainID helps close that gap by making access, action and data-exposure policy enforceable at runtime across the full agentic flow.

## References

- [Regulation \(EU\) 2024/1689 — consolidated EU AI Act](#) — Primary legal text, including Articles 9–17 and staged requirements.
- [European Commission — AI Act regulatory framework](#) — Current application timeline and implementation context.

#### Scope note

This document is a technical and product-capability mapping for customer discussion. Applicability depends on the organization's role in the AI value chain, use case, jurisdiction, risk classification, implementation architecture, and control design. Customers should validate legal, regulatory, audit, and certification conclusions with their qualified advisors.

