

■ COMPLIANCE FRAMEWORK REPORT · EU AI ACT · 1 OF 11 FRAMEWORKS

Evidence, not *assertion*.

Six EU AI Act articles assessed against this tenant's scan evidence — four evidenced, one partial, and one AIRM cannot evidence at all. It says so.

6

ARTICLES ASSESSED

4

EVIDENCED FROM THIS
SCAN

1

PARTIALLY EVIDENCED

1

OUTSIDE AIRM'S
EVIDENCE – STATED



What "evidenced" *means here*

This report maps scan evidence to EU AI Act obligations from the **deployer's** perspective — the organisation using AI systems, not building them. It is deliberately honest about scope: where the identity layer cannot evidence an obligation, the report says so instead of stretching.

● EVIDENCED

The scan produces direct, dated, reproducible evidence relevant to the obligation — an inventory, a log, a recorded oversight step.

● PARTIAL

The scan evidences part of the obligation; the remainder needs another control in your stack. The gap is named.

● CANNOT EVIDENCE

The obligation is a property of the AI model or its vendor, not of the identity layer. AIRM does not claim it — and tells you where to look instead.

Why honesty is the feature. Compliance packs that claim everything evidence nothing. An auditor who sees one obligation marked "AIRM cannot evidence this" trusts the other five findings — that is the point. Framework mapping shows findings are *relevant* to each obligation; it is not a certification.

EVIDENCE BASE FOR THIS REPORT

Read-only scan of 02 August 2026: 1,782 non-human identities inventoried, 234 AI applications identified and risk-assessed, ownership recorded or explicitly absent, 30-day Copilot activity audit, six priority findings surfaced for human decision.

THE OTHER TEN FRAMEWORKS

The same evidence maps to ISO/IEC 42001, NIST AI RMF, ISO/IEC 27001, DORA, UK CAF, MAS TRM, Essential Eight, CERT-In, BSI IT-Grundschutz and Cyber Essentials — one report per framework, generated from the same scan.

Six obligations, *assessed*

Art. 14 — Human oversight

● EVIDENCED

All 234 AI-related identities were surfaced for human decision; nothing was changed automatically. The six priority findings each carry a named decision (sanction / scope / revoke), and the oversight step is recorded with the scan date. Read-only by design — the human stays in the loop.

Art. 12 — Record-keeping

● EVIDENCED

AI identity activity is logged from the tenant's own audit trail: ~6,400 Copilot interactions over 30 days, ~55 distinct files accessed, permission grants and consent scopes recorded scan-over-scan. Records are dated, reproducible, and exportable (see the Identity Inventory Export).

Art. 26 — Obligations of deployers

● EVIDENCED

A complete inventory of AI systems in use (234, with origin and publisher), monitoring of their operation, and assignment of oversight: owners recorded or explicitly absent — 501 identities flagged as unowned, 61 of them critical, with remediation actions sequenced.

Art. 9 — Risk management system

● EVIDENCED

Continuous, repeatable risk identification for AI identities: blast-radius measurement, publisher trust, threat-intelligence matching (one watchlist hit surfaced), and scan-over-scan anomaly comparison — a standing risk-management loop for the AI identity layer.

Art. 10 — Data & data governance

● PARTIAL

What is evidenced: what data each AI system can reach (permission scope) and did reach (zero labelled-sensitive files read in 30 days). **The gap:** governance of training data and dataset quality inside each AI vendor's model — that sits with the vendor and your procurement diligence, not the identity layer.

Art. 15 — Accuracy, robustness & cybersecurity of the AI system

● AIRM CANNOT EVIDENCE THIS

Model accuracy and robustness are properties of the AI system itself, testable only by its provider. AIRM monitors the identity and access layer around the model — it does not test models, and this report will not pretend otherwise. Ask your AI vendors for their Art. 15 conformity documentation.

Three moves to *audit-ready*

The scan already evidences four of six obligations. These three actions convert the remaining exposure into documented, dated evidence.

- 1

Resolve the six priority findings ● CRITICAL

An open watchlist match and five unreviewed tenant-wide AI grants are exactly what an Art. 9 / Art. 26 review would flag. Decide each; the next scan records the closure.
- 2

Assign owners to the 61 critical unowned identities ● HIGH

Art. 26 oversight needs a name against every AI system. Owner assignment is recorded scan-over-scan — visible progress for an auditor.
- 3

Collect vendor Art. 15 documentation ● MEDIUM

For each sanctioned external AI app (post-review), request the provider's accuracy/robustness conformity documentation — the one obligation this report cannot evidence for you.

ASSESSMENT	STATUS	WHERE THE EVIDENCE LIVES
Art. 14 · Human oversight	● EVIDENCED	Priority findings register + recorded decisions (Executive Summary, pp. 5–6)
Art. 12 · Record-keeping	● EVIDENCED	30-day activity audit + Identity Inventory Export (CSV)
Art. 26 · Deployer obligations	● EVIDENCED	AI system inventory + ownership register (NHI Risk Report, pp. 3–4)
Art. 9 · Risk management	● EVIDENCED	Blast-radius scoring + threat-intel matching + anomaly engine (NHI Risk Report, p. 2)
Art. 10 · Data governance	● PARTIAL	Reach + exposure evidenced here; training-data governance → vendor diligence
Art. 15 · Accuracy & robustness	● CANNOT EVIDENCE	Provider conformity documentation — request from each AI vendor

ONE SCAN · ELEVEN FRAMEWORKS

The same dated evidence regenerates for every framework your auditors, board and insurer ask about — *without another questionnaire.*

Prepared using Sabiki AIRM — Agentic Identity Risk Management. Framework mapping evidences that findings are relevant to each obligation from the deployer perspective; it is not a certification, legal advice, or a conformity assessment under the EU AI Act, and a formal audit requires assessment beyond AI-identity monitoring. Findings reflect the read-only scan of 02 August 2026. This sample edition is fully de-identified: client name, domains, user identities and application-specific names are replaced, and all figures are adjusted by a 10–20% margin from an underlying production scan so that no value is attributable to any client; figures remain internally consistent. Tier-1 elements (NHI Risk Score™, movement deltas, reach counts) use illustrative sample values and render live once the corresponding portal fields are emitted. Sabiki Pte. Ltd. (UEN 202135394K), Singapore · sabikisecurity.com