

ISO 42001

The AI management system standard, explained for people who run companies.

by Nesibe Kırış Can · techletter.co

ISO 42001 does not make your AI trustworthy. It makes your AI governance auditable, and that is the point.

One warning: an ISO 42001 certificate is not solely EU AI Act compliance. It is evidence of governance, not a legal safe harbour.

KEY FACTS

What	ISO/IEC 42001:2023, the world's first certifiable AI Management System (AIMS) standard. Voluntary, audited by independent third parties, same management-system logic as ISO 27001. You can run an AIMS without ever being audited, and certify when a customer or market requires it
Who	Any organisation that develops, deploys, or uses AI. Any size, any sector, not only big tech. If you only buy and deploy vendor AI, you are still in scope.
Why early	Avoids last-minute compliance costs, meets procurement requirements, and signals readiness to customers, investors, and regulators.

WHERE IT SITS

NIST AI RMF UNITED STATES Voluntary risk vocabulary. A shared language, not a certification. VOCABULARY	EU AI Act EUROPEAN UNION Binding law with fines. Mandatory for AI placed on the EU market. LAW	ISO 42001 GLOBAL Voluntary standard, independently certifiable. The proof you hand over. CERTIFICATE
--	---	---

AND FOR AI AGENTS

Jan, 2026: Singapore's IMDA published the **Model AI Governance Framework for Agentic AI**, the first governance framework dedicated to AI agents. It answers the question the management-system standards leave open: how much autonomy, with how much oversight.

L1 HUMAN-LED Agent proposes, human operates.	L2 COLLABORATIVE Agent and human work together.	L3 SUPERVISED Agent operates, human approves.	L4 AUTONOMOUS Agent operates, human observes.
--	---	---	---

As autonomy scales, so must accountability.

Four governance dimensions: assess and bound risks upfront · make humans meaningfully accountable · implement technical controls · enable end-user responsibility.

THE STANDARDS AROUND IT					
42005 impact assessment	23894 AI risk	38507 governance	22989 terminology	23053 ML framework	
5338 AI system lifecycle	5259 data quality	TR 24027 bias	TR 24028 trustworthiness	TR 24368 ethics & society	

See also · [EU AI Act Cheat Sheet](#) See also · [Singapore IMDA Model AI Governance Framework Cheat Sheet](#)

02 · THE SYSTEM

What an AIMS actually is

The organisational structure of policies, processes, and designated roles that defines how AI is developed, deployed, and monitored. Not software. A way of running the company.

FIVE COMPONENTS

1

Documented scope

Which AI systems are in, and which are explicitly out.

2

Risk assessment

What can go wrong for the organisation. **Clause 6.1.2.**

Impact assessment

What can go wrong for the people and society the system touches. **Clause 6.1.4.**

|

Two different questions. The standard requires both. Defined in planning (6.1.2, 6.1.4); repeated at planned intervals in operation (8.2, 8.4).

3

Documented controls

For fairness, privacy, transparency, data quality, and accountability.

4

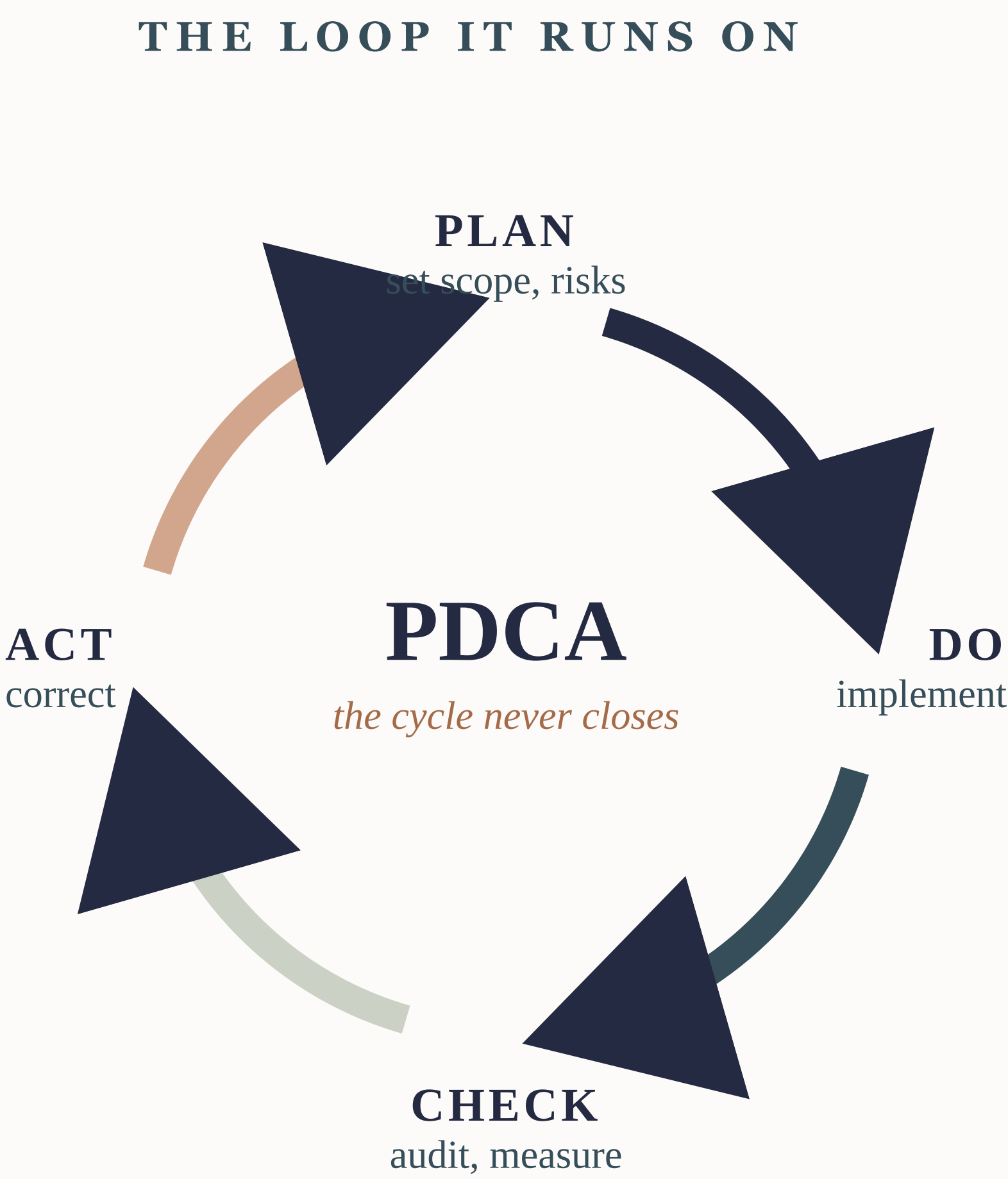
Trained, accountable staff

People with clearly assigned responsibilities, not job titles alone.

5

Measurement & correction

KPIs, internal audits, management reviews, and corrective actions.



WHAT IT LOOKS LIKE IN MONTH ONE

01

Inventory your AI systems

02

Name the owner

03

Write the risk criteria

04

Start the document trail

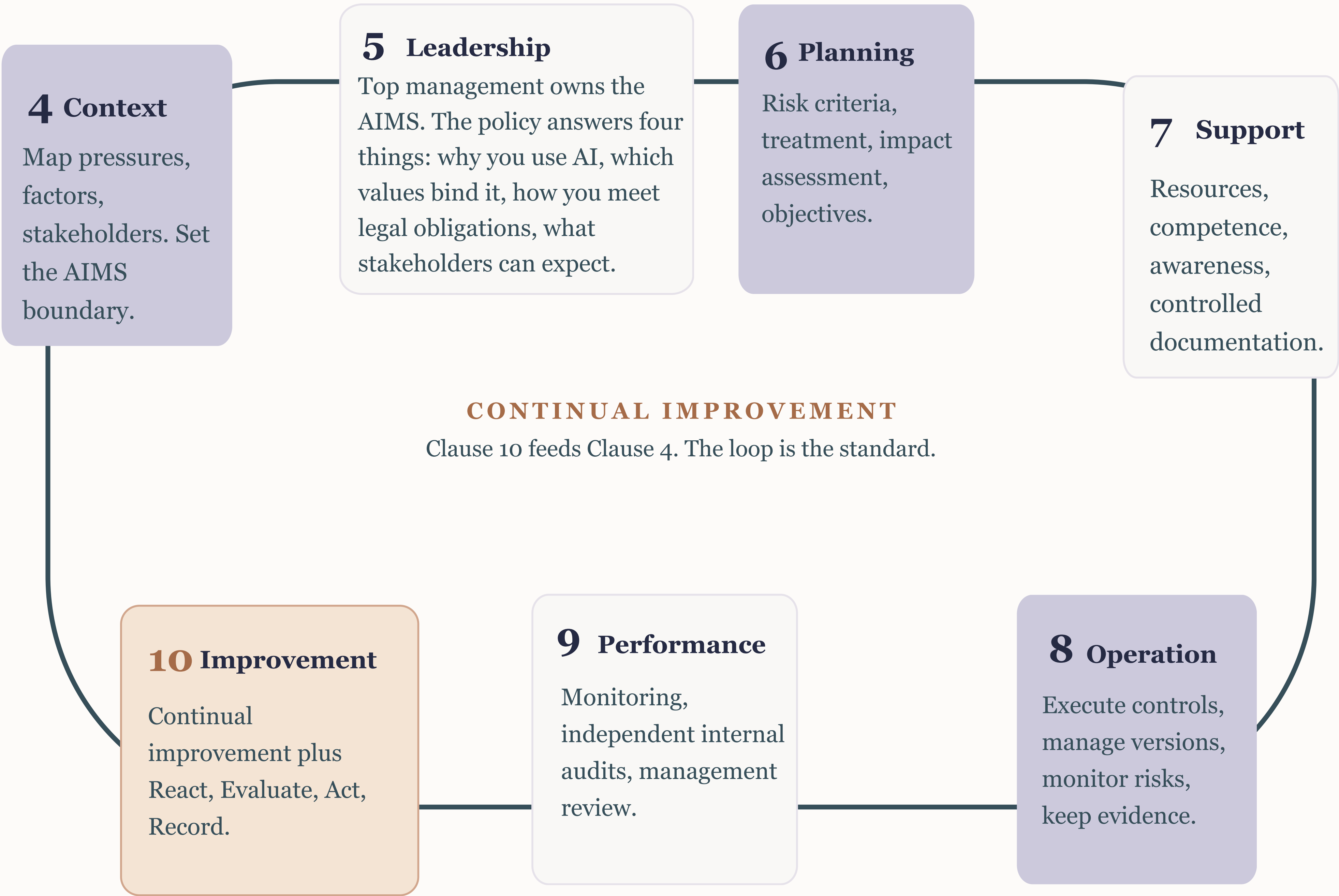
🔗 If you already run ISO 27001, the AIMS reuses the same machinery: the standard shares the harmonized structure of ISO/IEC 27001, 27701, and ISO 9001. **Annex D's argument** is that AI governance should not run as a silo; security, privacy, and quality are managed across the whole system.

An AIMS is an operating system, not a project.

03 · THE STANDARD

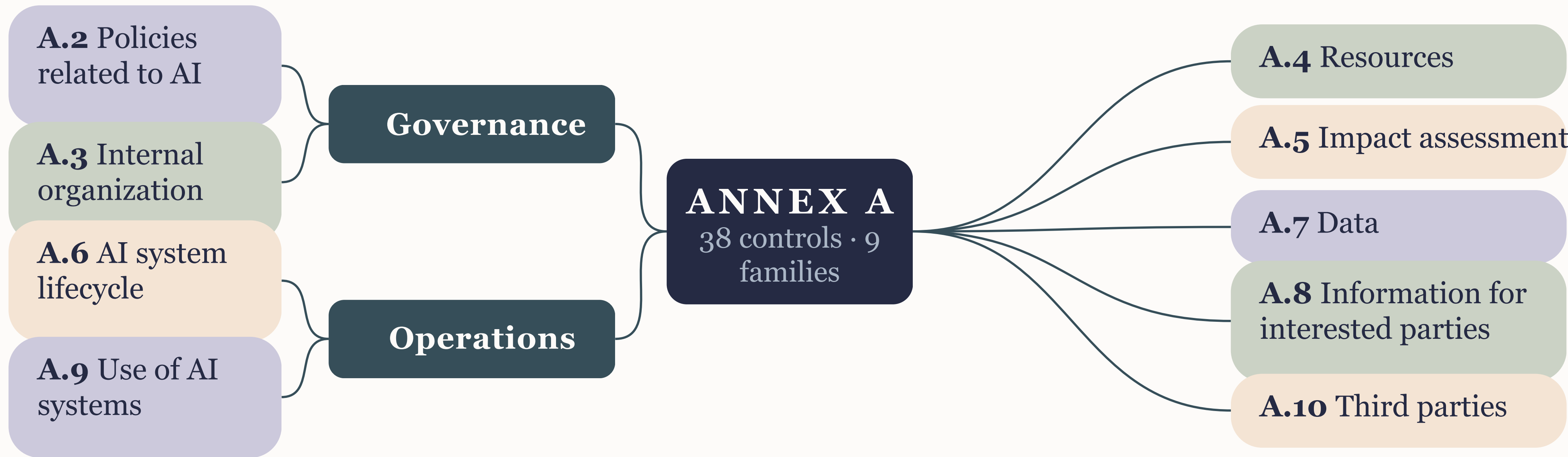
The clause map, 4 to 10

Seven clauses carry the requirements, and they run as a loop: Clause 10 feeds back into Clause 4.



THE CONTROLS BEHIND THE CLAUSES

Clauses 4-10 define the management system. Annex A defines the concrete controls you select from, in 9 families.



9 families, 10 objectives: A.6 carries two objective groups, management guidance (A.6.1) and the system life cycle itself (A.6.2).

Statement of Applicability (Clause 6.1.3): the document declaring which Annex A controls apply to you and why the rest do not. Auditors read it first. Annex B gives implementation guidance for every control; you are not left to interpret them alone.

04 · THE DECISIONS

Seven founder decisions

One per clause. These are the calls only leadership can make.

1

📍

CLAUSE 4

Decide where your AIMS stops. A defined boundary is legitimate; an undefined scope is not.

The standard prescribes no timeline or budget. Effort scales with scope, which makes this **the cost lever**.

2

👑

CLAUSE 5

Accountability stays at the top. You can delegate implementation to your CTO, not ownership.

3

📋

CLAUSE 6

Write down the risk you accept. Residual risk is documented and formally accepted by management, not silently absorbed by the team.

4

📖

CLAUSE 7

Budget for competence, not just tooling. Training, documentation, and communication are audit evidence.

5

🔌

CLAUSE 8

Kill switch and rollback exist before launch, not after the incident. Every significant change gets an assessment and a decision trail.

6

🔗

CLAUSE 9

The person who audits cannot be the person who builds. Independence is structural, not a formality.

7

🔄

CLAUSE 10

Failures are inputs. An issue that repeats is a governance failure, not a technical one.

WHAT YOU CHOOSE TO OPTIMIZE FOR

ANNEX C

Annex C is the standard's menu, explicitly non-exhaustive: the objectives you can manage risk against, and where risk comes from. Deciding which of these matter to you is **Decision 3** in practice.

OBJECTIVES

accountability

AI expertise

availability & quality of training and test data

fairness

maintainability

privacy

robustness

safety

security

transparency & explainability

environmental impact

RISK SOURCES

complexity of environment

lack of transparency & explainability

level of automation

machine learning risks, incl. data quality and poisoning

system hardware issues

technology readiness

life cycle issues

Fuller treatment in ISO/IEC 23894, AI risk management.

05 · THE PEOPLE & THE PATH

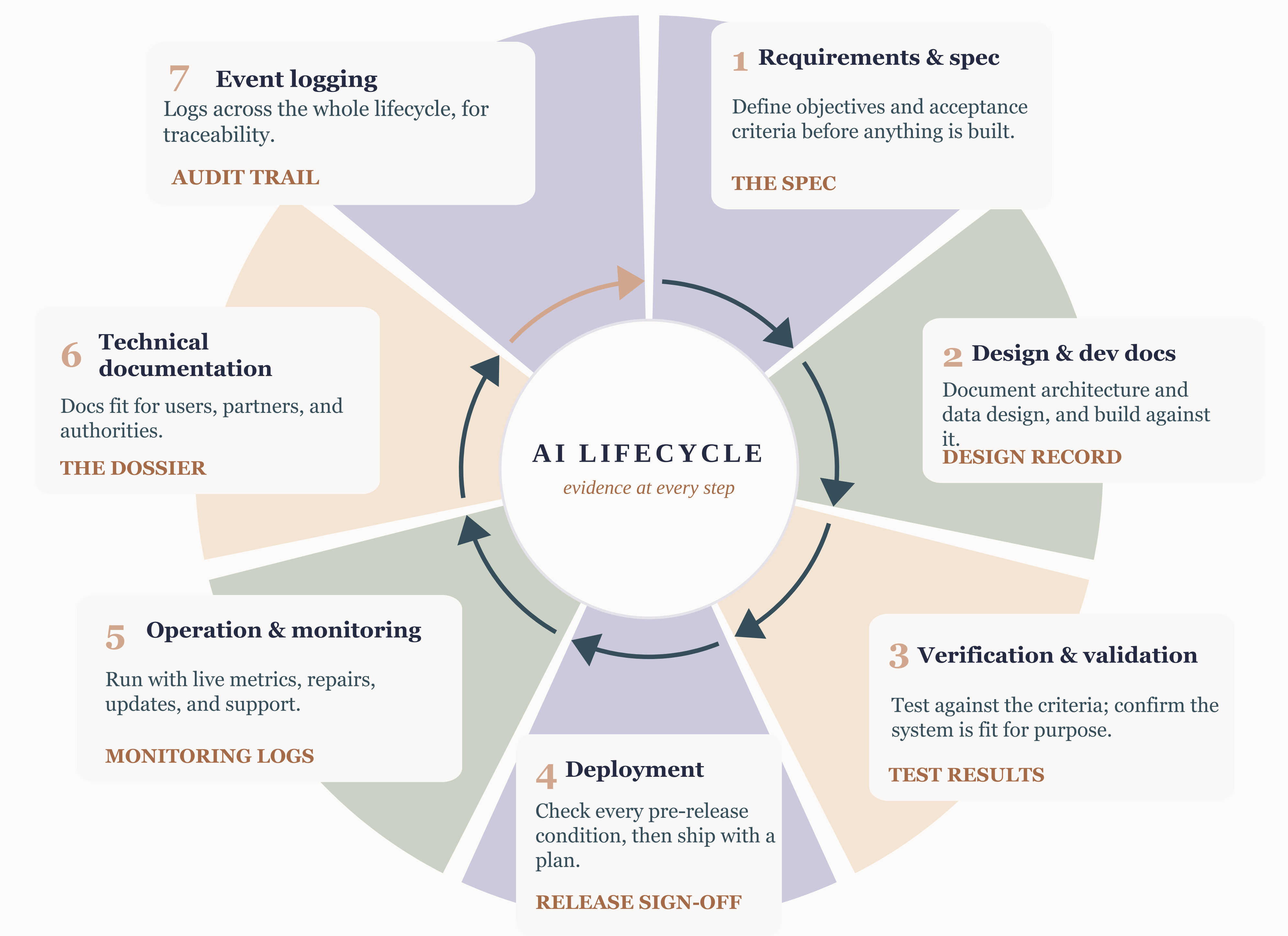
Roles and the lifecycle

MINIMAL VIABLE ROLE SET

ROLE FAMILY	WHAT THEY OWN
Executive / Governance	Accountability, the AI policy, and the resources behind it.
Data	Quality, provenance, and protection of what feeds the model.
Design / Engineering	Building the system and its oversight mechanisms.
Deployment / Operation	Running it, monitoring it, and responding to incidents.
Auditing / Monitoring	Checking the whole thing independently.

Three separations are non-negotiable, even when one person holds several roles: executive accountability at the top, auditor independence from daily operations, and a named owner for human oversight of consequential decisions.

THE LIFECYCLE · EVIDENCE AT EVERY STEP



And at the end of life: retiring a system is a governed event too. Data disposed per retention rules, users notified, records archived.

06 · IN PRACTICE

A worked example, and where it breaks

A FICTIONAL CASE, A REAL PATTERN

FICTIONAL CASE STUDY, FROM THE SOURCE REPORT

X Corporation · AI HR pre-screening (scope: ranking only, humans keep the decision)

●
Deploy
ranking live

— ●
78%
fairness alert

— ●
Independent audit
thresholds too high

— ●
Retrained
thresholds fixed

— ●
82→91%
within thresholds

The system worked because monitoring, auditor independence, and documentation existed before the failure, not after it.

WHERE IMPLEMENTATIONS BREAK



Certification treated as a project with an end date, not a management system that keeps running.



Policies without named owners and deadlines. A policy nobody owns is a document, not a control.



Auditors reporting to the people they audit. Independence collapses and the audit is theatre.



Residual risk never documented or formally accepted, so nobody owns the exposure.



Vendor AI tools bought with no vendor assessment or contract-level risk allocation. Their risk becomes yours, silently.



A risk register that names no specific system. Generic risks produce generic controls, and auditors notice.

07 · TAKE IT AWAY

The founder's checklist

- 4 · CONTEXT

 - Every AI system in use is inventoried, including third-party and vendor AI
 - AIMS boundary documented, in and out
 - Stakeholders and legal requirements mapped
- 5 · LEADERSHIP

 - AI policy approved and communicated
 - Named executive owner
 - Roles and authorities assigned in writing
- 6 · PLANNING

 - Risk criteria defined, Low / Medium / High
 - Treatment plan approved
 - Residual risk accepted in writing by management
 - Measurable objectives with owners and deadlines
- 7 · SUPPORT

 - Training delivered, competence periodically evaluated
 - Documentation under version control
 - Internal and external communication defined: what, to whom, when
- 8 · OPERATION

 - Controls implemented and evidenced
 - Kill switch and rollback tested
 - Every change logged: who, why, decision, outcome
 - Incident response path defined and tested, with a named owner
- 9 · PERFORMANCE

 - Monitoring with alert thresholds live
 - Independent internal audit scheduled
 - Management review at planned intervals
- 10 · IMPROVEMENT

 - Nonconformity procedure: React, Evaluate, Act, Record
 - Corrective actions documented and verified

EVIDENCE · WHAT YOU SHOW, NOT WHAT YOU INTEND

- AI inventory with named owners
 - Monitoring reports and alerts
 - Statement of Applicability kept current
- Risk register with treatment status
 - Human review and override records
- Approved change logs, models and data
 - Vendor assessments

Audits verify execution, not intentions.

Based on the ISO 42001 Starter Guide (HUX AI, 2025), co-authored by Nesibe Kırış Can with Burçin Kızılcıklı, Ege Uğur Amasya, Hayriye Anıl, İdil Kula, and Onur Pişirir. Licensed CC BY-NC 4.0.
Primary standard: ISO/IEC 42001:2023 · [iso.org/standard/42001](https://www.iso.org/standard/42001)

TechLetter covers AI governance, policy, and the incentives underneath them.
Weekly, for people who make these decisions

- Email: me@nesibekiris.com
- LinkedIn: linkedin.com/in/nesibe-kiris
- Twitter/X: [@nesibekiris](https://twitter.com/nesibekiris)

