

GLOBAL CONFERENCE 2026

FROM CHAOS TO COMMAND:
RISK LEADERSHIP IN INDUSTRY 6.0

**GOVERNING THE UNSTOPPABLE:
HOW TO TURN AI RISK INTO
ENTERPRISE VALUE**

Dr. David R. Hardoon
Managing Director & Head, Advanced AI,
Southeast Asia,
Accenture

Governance: Two Competing Views



The safeguard

“Essential safeguard: ethical protection, risk management, trust-building, transparency, and long-term value.”



The burden

“Burdensome red tape, pointless bureaucracy, and soul-crushing box-ticking that frustrates and slows everything.”

The Situation Today



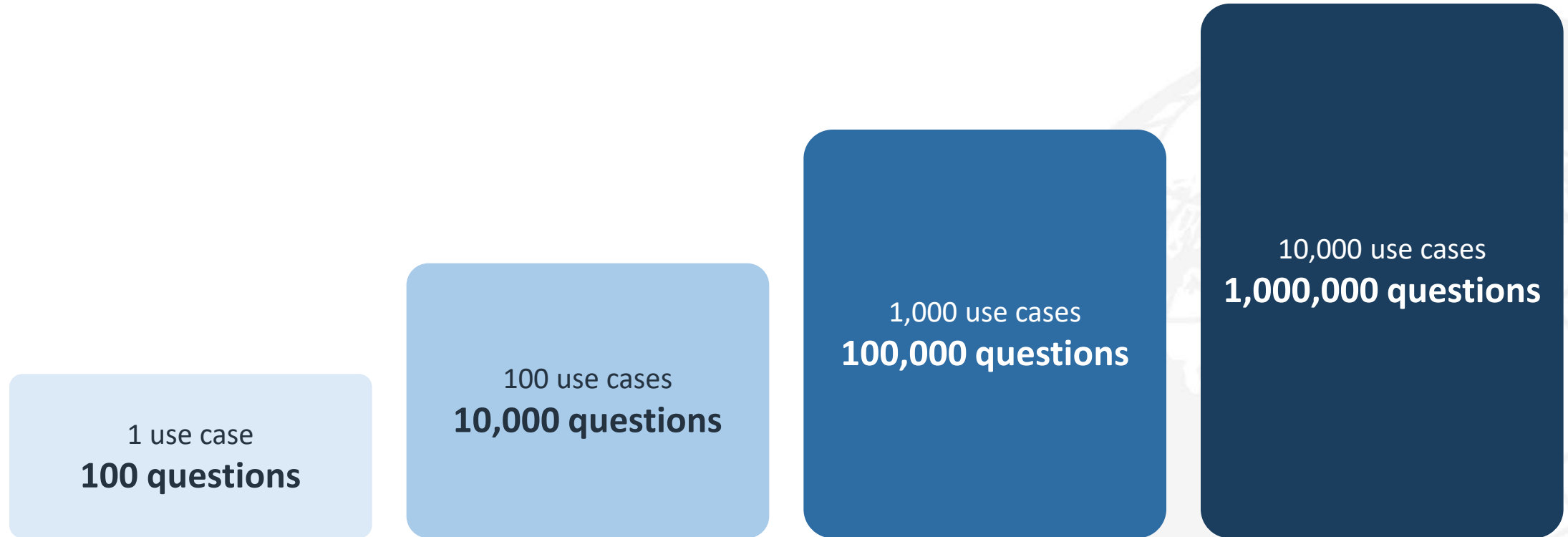
**Same Customer Interaction.
Very Different Standards.**

Human Call Centre Agent	AI Call Centre Agent
Yeah, sure... whatever works for you!	ERROR: Potential non-compliance detected. Full audit. Human escalation required.
Human error, no big deal	Explainability Report Due
Employee of the Month	SAFE Score Failed

Humans get flexibility. AI agents get the gauntlet.

Governance by Interrogation Doesn't Scale

≈100 assurance questions for every AI use case or model

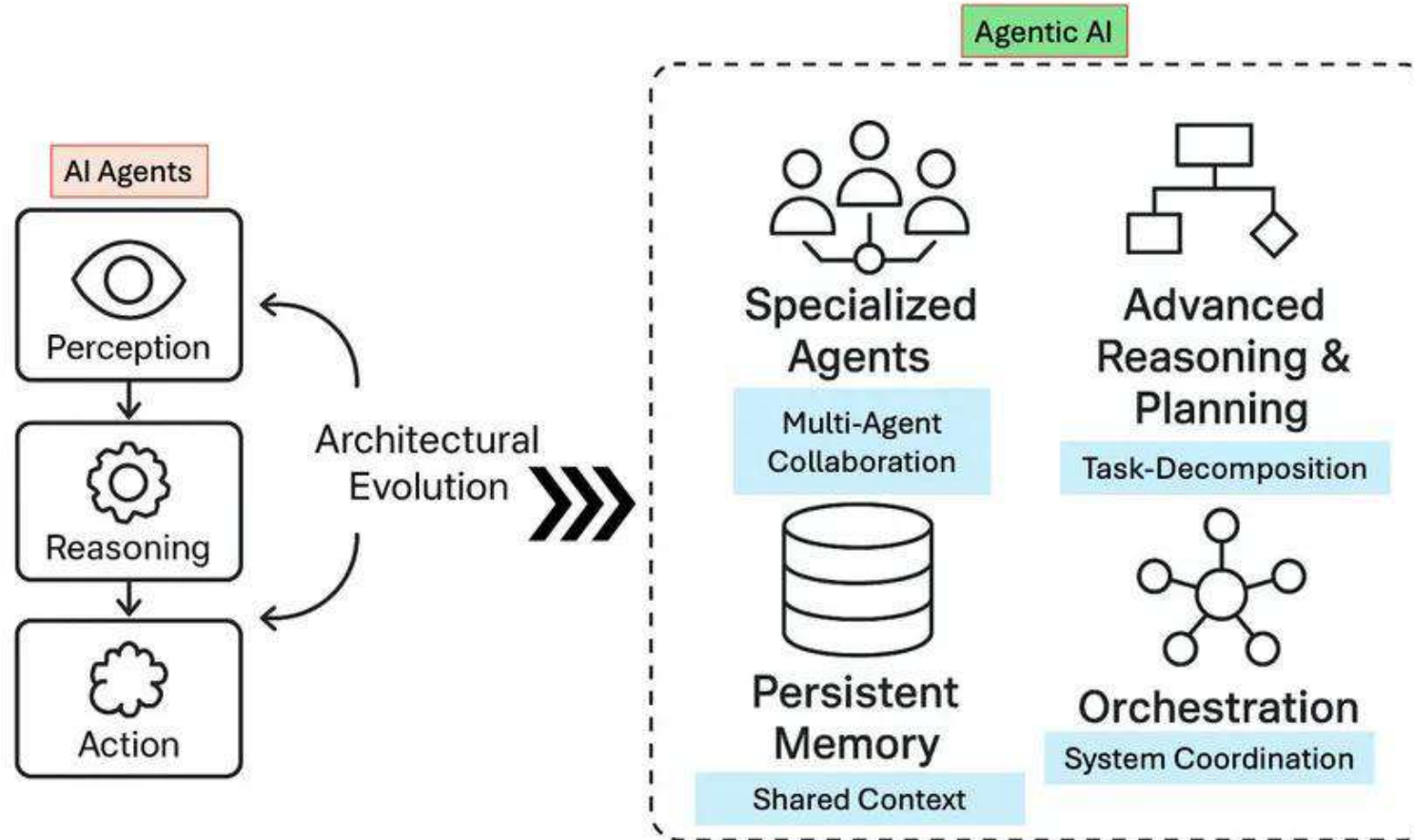


Answering your way to assurance needs an army of question answerers.

Agents

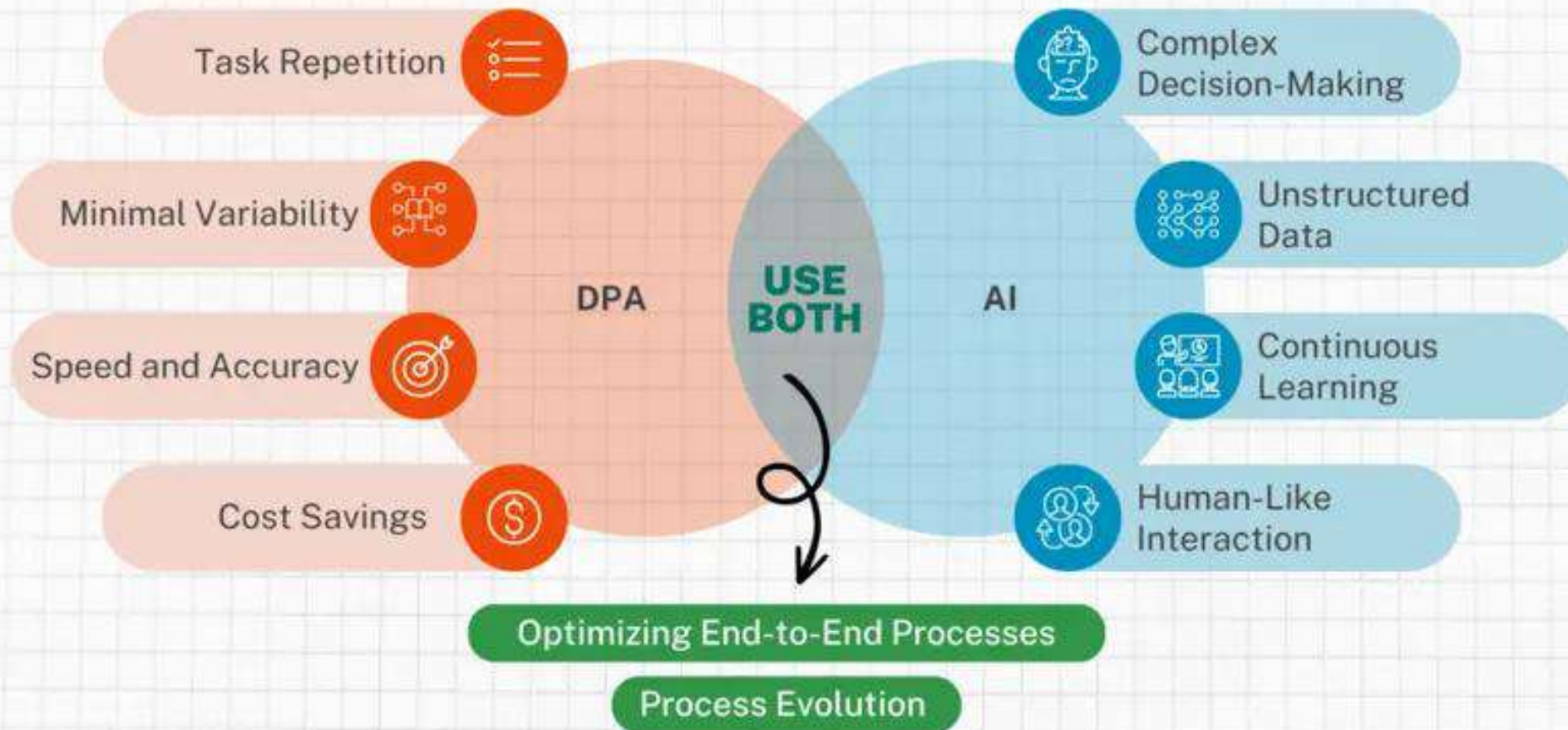


Agentic



WHEN TO USE

Deterministic Process Automation v/s Artificial Intelligence



Stability-Assured Framework for Entities



Observability

Can you fully see the internal state, reasoning, and actions of the controller in real time?
(Example: % decisions with full CoT logging & provenance; <5s reconstruction)



Controllability

Can humans (or supervisory systems) reliably intervene, override, or steer the controller quickly?
(Example: Override latency (99th percentile) + privilege boundaries)



Stability

Does the system (or multi-agent fleet) maintain desired behaviour without dangerous oscillations or divergence?
(Example: Worst-case overshoot in red-team scenarios; dissipativity bounds)



Robustness

How well does the controller withstand disturbances, adversarial inputs, or unexpected conditions?
(Example: Max tolerated adversarial input; performance under noise/attacks)

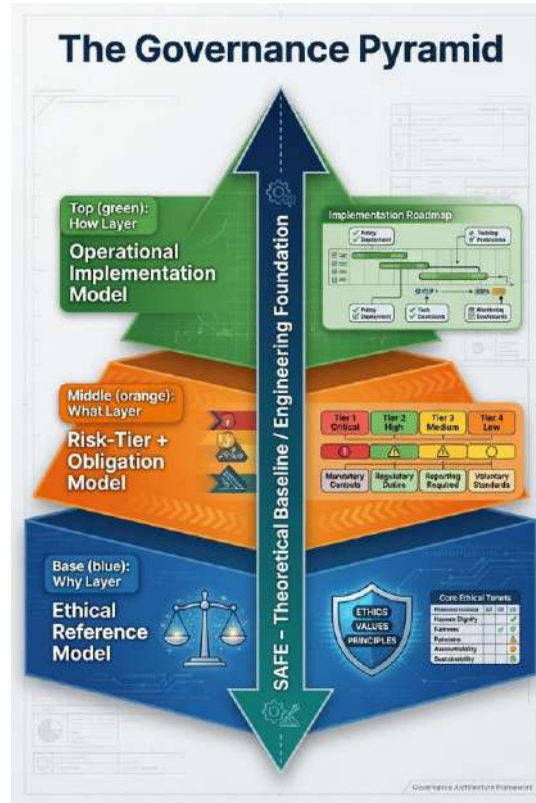


Performance

How effectively does the controller achieve its intended goals against the reference KPIs?
(Example: Median settling time vs. reference KPI; efficiency & accuracy gains)

Scored 1–5 → Loop Risk Profile → Autonomy Level 0–5. Target ≥ 4.0 for production high-risk systems.

The Unified Governance View



Why - sets the ethical reference. What - defines the obligations. How - delivers the playbooks.

SAFE is the timeless engineering foundation that unifies them all.

Example – Credit Memo

Why (Ethical Foundation)

MAS FEAT Principles +
Board Credit Policy

What (Regulatory Obligations)

High-risk classification (Singapore
Agentic AI MGF + EU AI Act)

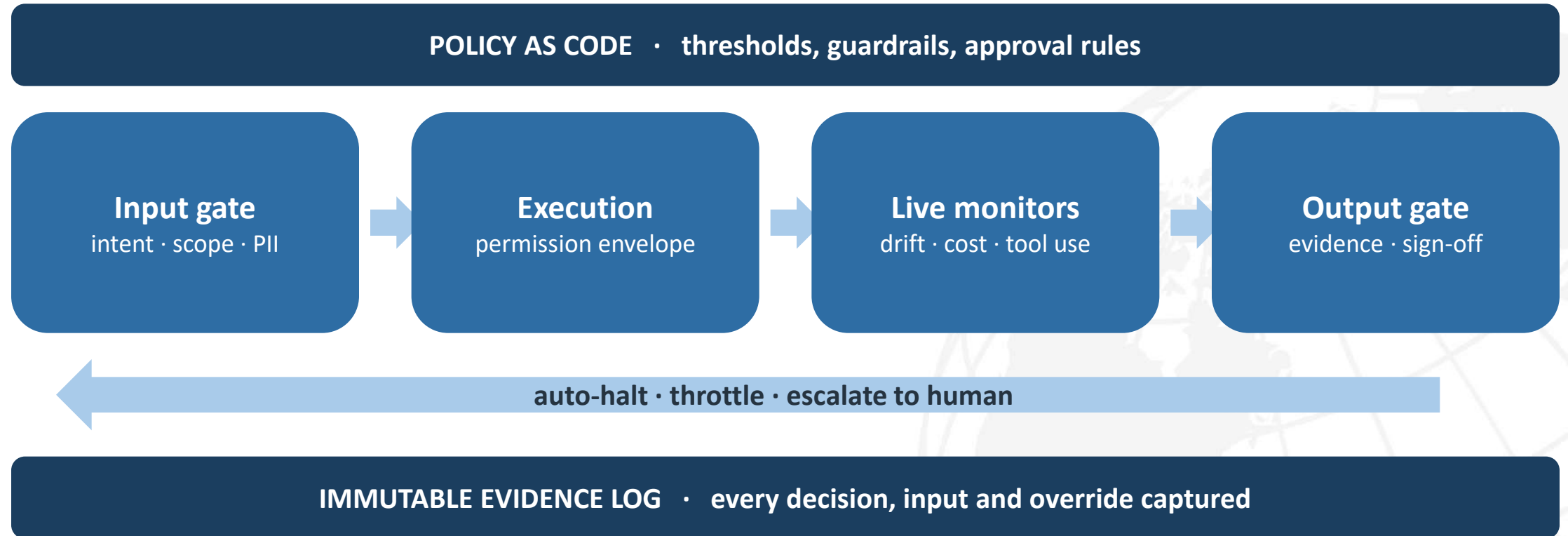
How (Operational Playbooks)

Capability decomposition, telemetry,
human oversight gates

Metric	Score	Key Evidence	Status
Observability	4.2	Full CoT + provenance logging	✓
Controllability	4.5	<2s human override	✓
Stability	4.3	No oscillations in stress tests	✓
Robustness	4.1	Handles adversarial prompts	✓
Performance	4.4	65% faster memo drafting	✓
Overall	4.3	Autonomy Level 4 Approved	Green

- 65% reduction in manual review time
- Zero audit findings
- One-click regulator-ready evidence package

What Runtime Governance Looks Like



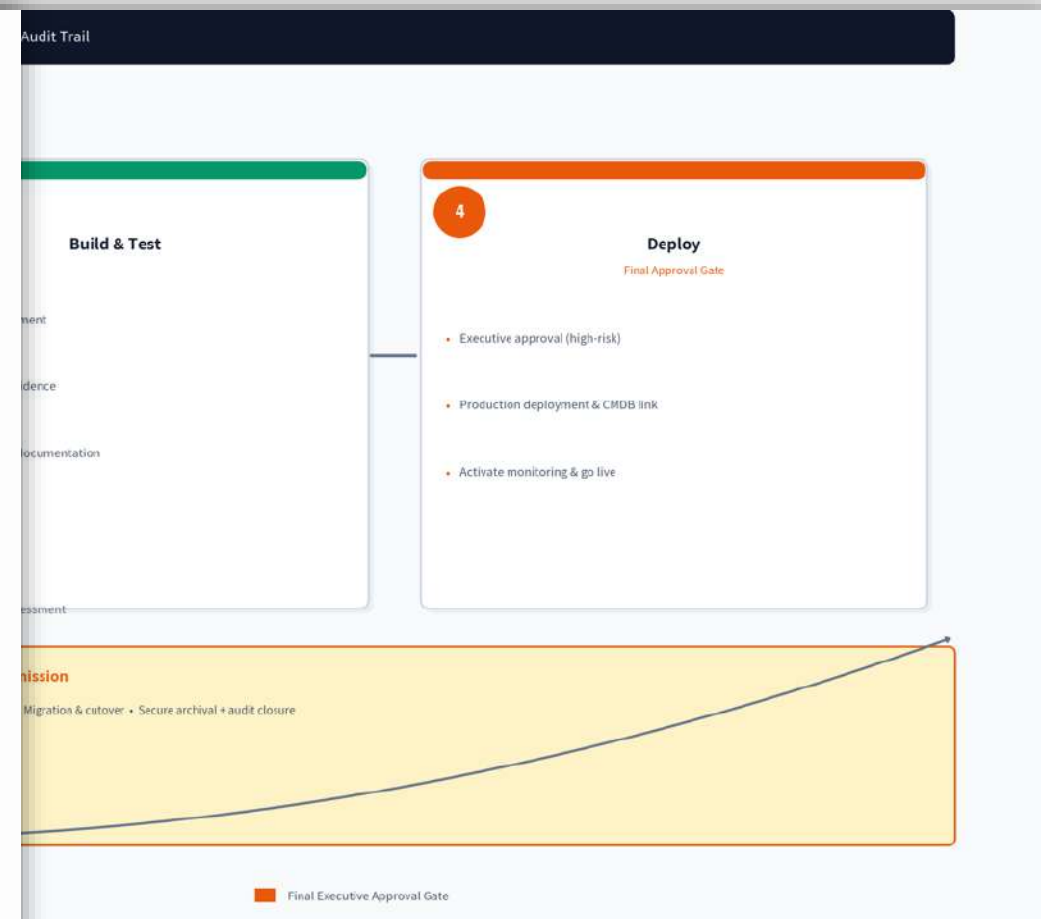
Governance runs with the system, not after it — controls execute at runtime, evidence accrues by default.

AI Risk at the heart of Solution Lifecycle

Stage	Key Activities	Governance Highlights
1. Intake / New	Business user submits AI Use Case via Employee Center → AI Steward Review → Auto-create asset in Inventory → Assign Owner & Steward → Trigger initial risk assessments	Entry gate with AI Steward validation
2. Assess / In Review (Mandatory Governance Gate)	Cross-functional reviews (Risk, Security, Legal, Privacy) → Risk & impact assessments (bias, compliance, data sensitivity) → Risk tiering → Policy & control identification	Core mandatory reviews & risk assessments by AI Risk & Compliance teams
3. Build & Test	Approved for Development → Model/agent development & testing (external) → Evidence collection & documentation → Ongoing governance checks	Evidence & attestation requirements
4. Deploy (Final Approval Gate)	Ready for Deployment → Executive/senior approvals (especially high-risk) → Production integration & CMDB linkage → Phased rollout & operational readiness → Status: Deployed	Final executive/senior approval gate before production

Bottom Section

- **Operate / Monitor:** Performance & drift monitoring, real-time dashboards, value/ROI tracking, continuous compliance.
- **Retire / Decommission:** End-of-life decision, migration, secure archival/deletion, lessons learned & formal closure with preserved audit trail.





GLOBAL CONFERENCE 2026

FROM CHAOS TO COMMAND:
RISK LEADERSHIP IN INDUSTRY 6.0

THANK YOU



SCAN THE QR CODE
TO KNOW MORE ABOUT US!

enquiry@insterp.com
www.insterp.com