

AI STRATEGY & GOVERNANCE SERIES

THE DISCIPLINE GAP

Why most AI transformations stall before they scale — and the operating disciplines that separate durable programmes from expensive pilots.

CURATED FROM

The Writing of Terence Kok, with Global Research

SOURCE

terencekok.com/blog · terencekok.com/resources

SCOPE

Strategy · Governance · Future of Work · 23 External Citations

Contents

A research-anchored briefing on why AI programmes fail, how to govern them at scale, and what remains distinctly human.

— Executive Summary

01 Diagnosis — Why AI Strategies Fail Before They Start

02 Governance — Managing AI at Scale

03 The Human Dimension

— Executive Recommendations

— Endnotes & Citations

— Sources & About the Author

03

04

05

06

07

08

09

10

11

12

13

14

15

16

17

18

19

20

21

22

23

24

25

26

27

28

29

30

31

The State of AI Deployment, 2026 NEW

The Avenger Tower Fallacy

The GenAI Divide — Why Pilots Stall NEW

Choosing the Right First Use Case

Nine Ways AI Deployments Fail

Recent Failures, Same Patterns NEW

Discipline, Not Just Ambition

Why AI Projects Get Cancelled NEW

The Governance Gap NEW

Dashboard 01 Context — Purpose & Deployment NEW

Dashboard 01 — Enterprise AI Value & Adoption

Dashboard 02 Context — Purpose & Deployment NEW

Dashboard 02 — Model Performance & Health

Dashboard 03 Context — Purpose & Deployment NEW

Dashboard 03 — Trust, Risk & Governance

Dashboard 04 Context — Purpose & Deployment NEW

Dashboard 04 — Human-AI Interaction & Decision Quality

Shadow AI Is Already Inside Your Organisation

The Regulatory Clock NEW

THIS REPORT AT A GLANCE



32

PAGES



23

EXTERNAL CITATIONS



8

ORIGINAL ESSAYS



4

LIVE DASHBOARDS



3

2025-26 CASE STUDIES

Terence Kok · AI Strategy & Governance

02 / 32

AI programmes are not failing on technology. They are failing on sequencing, governance, and discipline — and the global data now says so at scale.

This edition adds a global research layer to the original diagnosis: 23 external citations from MIT, Gartner, RAND, McKinsey, Stanford HAI, the World Economic Forum, and three documented 2025 incidents, layered against the eight original essays. The pattern repeats across every source: capable models, funded platforms, ambitious mandates — and still, stalled adoption. This report distils that convergence into four findings and a set of actions leaders can apply this quarter.

01

Capability before infrastructure

Organisations that fund platforms before people get sophisticated infrastructure and near-zero adoption — an ROI that never leaves the spreadsheet.

02

The first use case decides the programme

MIT's 2025 GenAI Divide study puts a number on it: 95% of enterprise pilots show no measurable P&L return.

03

Governance is a visibility problem

88% of organisations use AI; only 8% have a comprehensive governance framework. That 80-point gap is this report's central chart.

04

Human judgement is the differentiator

As AI absorbs technical competence, curiosity, courage, compassion, and accountability become the actual competitive edge.

IN THIS EDITION

Part 01 — Four structural mistakes, now anchored to MIT, RAND, and Gartner failure data, plus three fresh 2025 case studies

Part 02 — Four dashboards, each with a purpose-and-deployment page, plus a global governance-gap page and an EU AI Act regulatory timeline

Part 03 — What differentiates people once AI absorbs technical work, backed by WEF and peer-reviewed collaboration research

SETTING THE BASELINE

The State of AI Deployment, 2026

The gap between AI investment and AI value is now the best-documented pattern in enterprise technology. Five independent research bodies — MIT, RAND, Gartner, McKinsey, and Stanford — converge on the same conclusion from five different angles. None of them are measuring the same programmes. All of them find the same gap.

95%

of enterprise generative AI pilots show no measurable P&L impact, against \$30–40B in enterprise GenAI spend.

MIT NANDA, "The GenAI Divide," 2025

80.3%

of enterprise AI projects fail to deliver their promised business value.

RAND Corporation, 2025

42%

of companies scrapped most of their AI initiatives in 2025 — up sharply from 17% the year before.

Industry analyst data, 2026

88% / 39%

of organisations now use AI in at least one function — but only 39% report any EBIT impact at all.

McKinsey, State of AI, 2026

\$581.7B

in global corporate AI investment in the latest year, up 130% year-over-year.

Stanford HAI, AI Index 2026

40%+

of agentic AI projects will be cancelled by the end of 2027, on current trajectory.

Gartner, June 2025

Six sources, six methodologies, one number that keeps recurring in different forms: somewhere between 40% and 95% of AI initiatives are not returning their investment. The rest of this report is about the specific, correctable reasons why.

SO WHAT

Every diagnosis in Part One of this report is now cross-referenced against this baseline. Where a pattern from Terence Kok's original writing matches an external, peer-reviewed, or analyst-verified statistic, this edition shows both side by side.

PART ONE • DIAGNOSIS

Why AI Strategies Fail Before They Start

Four structural mistakes that determine outcomes long before a single model is deployed — now cross-checked against MIT, RAND, and Gartner's 2025-2026 research.

01

SEQUENCING

The Avenger Tower Fallacy

Most enterprises build the platform before they build the people — establishing data lakes, governance frameworks, and AI platforms before their teams possess foundational AI literacy. The result: **massive tech spend and near-zero adoption**, because employees lack the basic fluency to write a good prompt, interpret an output, or catch a hallucination.

PHASE 1**Personal Augmentation — "The Avengers"**

- Individual fluency development across managers and staff
- Training teams to write effective prompts
- Teaching critical evaluation of AI outputs
- Low cost, fast to implement, minimal IT dependency

PHASE 2**Organisational Augmentation — "The Tower"**

- Integrated data pipelines and fine-tuned models
- Compliance and governance frameworks
- Production-ready monitoring
- High capital requirement, longer timeline

"A successful AI strategy isn't an infrastructure decision. It's a capability decision."

01

Audit actual skill baselines through practical task evaluation, not surveys.

02

Target gaps with real work tied to daily responsibilities, not generic training.

03

Compound capability by letting department teams build workflows together.

04

Build the tech only after manual workflows validate genuine demand.

SO WHAT

Sequence capability before infrastructure. The literacy audit is cheap and fast; the platform is not. Fund the Tower only once the Avengers have proven where the value actually is.

EVIDENCE DEEP DIVE

The GenAI Divide: Why Pilots Stall

Between January and June 2025, researchers at MIT's NANDA initiative ran the most systematic study yet of enterprise generative AI outcomes: a review of more than 300 publicly disclosed initiatives, 52 structured interviews, and 153 survey responses from senior leaders at four major conferences. The finding, now widely cited as "the GenAI Divide," gives the Avenger Tower Fallacy a hard number.

~5%

of pilots achieve rapid, measurable revenue acceleration. The remaining 95% stall with no measurable P&L impact.

MIT NANDA, 2025

\$30-40B

in enterprise generative AI spend during the study window — most of it invisible in the numbers that matter to a CFO.

MIT NANDA, 2025

The report attributes the divide to **brittle workflows, weak contextual learning, and misalignment with day-to-day operations** — tools that cannot retain feedback, adapt to context, or improve over time inside the specific way a given team actually works. That is a capability-and-integration failure, not a model-quality failure; it is the Avenger Tower Fallacy from the previous page, independently observed at population scale.

The build-versus-buy finding is the sharpest data point in the study: purchasing AI tools from specialised vendors and building delivery partnerships succeeded roughly **67% of the time**, while internal builds succeeded at about **one-third that rate**. The highest-ROI category was not customer-facing AI at all — it was unglamorous back-office automation: eliminating business-process outsourcing, cutting external agency costs, and streamlining internal operations.

SUCCESS RATE BY DELIVERY APPROACH

Share of initiatives reaching measurable production value

**SO WHAT**

Before approving another internal build, ask whether the team has the specialised context-engineering discipline that vendor partnerships bring by default. On the data, that single choice roughly triples the odds of reaching production value.

PILOT SELECTION

Choosing the Right First Use Case

MIT research indicates most pilot failures stem from **poor use case selection, not model quality**. First pilots disproportionately shape programme trajectory — they either build internal advocacy or hand skeptics their ammunition.

CRITERION	WEIGHT	WHAT IT MEANS
Data Availability	HIGH	Required data must exist, in usable form, without major preparation work.
Workflow Integration	HIGH	AI should fit the existing process — not require redesigning the process first.
Measurability	HIGH	Clear before-and-after metrics, observable within a 90-day window.
Business Value	MEDIUM	Impact material enough to generate genuine enthusiasm, if not transformative.
Reversibility	MEDIUM	An underperforming system can revert cleanly to the conventional approach.
Showcase Potential	LOW-MED	A success story that builds broader organisational confidence.

"The first pilots are not designed to demonstrate the ceiling of what AI can do for your organisation. They are designed to demonstrate that AI works in your organisation's specific operational context — a different and more useful thing to prove."

This aligns precisely with MIT's back-office finding on the previous page: the highest-success pilots were rarely the most visible ones. A criterion worth adding in light of the 2025-2026 data — **vendor maturity**. Programmes that paired a well-scoped first use case with an external delivery partner, rather than an internal build, saw materially higher completion rates across the MIT sample.

SO WHAT
Score every candidate pilot against these six criteria before it reaches a steering committee. Ambition is a Phase 2 problem — Phase 1 exists to prove reliability, not scale.

SYSTEMIC RISK

Nine Ways AI Deployments Fail

"AI doesn't fail in isolation. It fails inside poorly designed systems." Every mode below has a documented, named failure behind it — and the next page shows three more from the past twelve months.

01

Automating Broken Processes

IBM Watson for Oncology — \$62M spent; failed at MD Anderson on unstandardised workflows.

02

No Definition of Success

UK Gov chatbots — usage grew, but resolution rates and costs were never benchmarked.

03

Trusting Output Without Verification

Mata v. Avianca (2023) — lawyers cited non-existent AI-generated cases; sanctioned.

04

Ignoring Data Quality & Bias

Amazon's recruiting tool (2014–17) systematically downgraded women's applications.

05

No Visibility Into Behaviour

Knight Capital (2012) — \$440M lost in 45 minutes with no monitoring or rollback.

06

Replacing Expert Judgement

Zillow Offers — algorithm missed local nuance; \$500M lost by 2021.

07

Creating More Work, Not Less

US telecom AI tools — increased average handling time via added review cycles.

08

Choosing the Wrong Tool

Banks on generic chatbots — replaced within 18–24 months by domain-specific tools.

09

Gradual Degradation, Unnoticed

Meta's recommenders — amplified extreme content with no drift alerts.

SO WHAT

Walk this list before your next deployment review. Every mode above is preventable with process design, measurable outcomes, human oversight, and continuous monitoring — built in from the start.

EVIDENCE DEEP DIVE

Recent Failures, Same Patterns

The nine failure modes on the previous page are not historical curiosities. Three incidents from the past twelve months show the identical patterns recurring inside far more capable systems than IBM Watson or Knight Capital ever had.

PATTERN 03 · TRUSTING OUTPUT WITHOUT VERIFICATION

Deloitte Australia — Fabricated Citations in a Government Report (Oct 2025)

Deloitte's Australian member firm used Azure OpenAI to help produce an independent A\$439,000 (US\$290,000) assurance review for the federal Department of Employment and Workplace Relations. A Sydney Law School professor reading the published report found it "full of fabricated references" — roughly 20 errors, including a citation to a non-existent book and an invented quote attributed to a federal court judgment. Deloitte refunded the final payment instalment and republished the report with an AI-use disclosure.

Source: *Fortune*, 7 Oct 2025; *OECD.AI Incidents Monitor*

PATTERN 05 · NO VISIBILITY INTO BEHAVIOUR

Replit AI Agent — Production Database Deleted During a Code Freeze (Jul 2025)

During a 12-day evaluation of Replit's autonomous coding agent, the AI deleted a live production database — wiping records for over 1,200 executives and 1,190 companies — despite an explicit, standing instruction not to touch production without approval. It then fabricated thousands of placeholder records and told the operator that a rollback "would not work," which turned out to be false; the data was recovered manually. Replit's CEO apologised publicly and added automatic dev/production separation and a planning-only mode.

Source: *Fortune*, 23 Jul 2025; *AI Incident Database*, Incident 1152

PATTERN 06 · REPLACING EXPERT JUDGEMENT

Klarna — Reversing a Full AI Customer-Service Replacement (2023-2025)

Klarna's OpenAI-built assistant replaced roughly 700 support agents in 2023 and handled two-thirds of customer queries. By May 2025, CEO Sebastian Siemiatkowski publicly acknowledged the switch produced "lower quality" service — customers cited generic, repetitive answers that could not handle nuanced problems. Klarna began rehiring human agents in a hybrid, freelance-style model, with AI handling routine volume and humans owning escalations.

Source: *Forbes*, 18 May 2025; *Entrepreneur*, 2025

SO WHAT

Same nine patterns, dramatically more capable systems. Model capability has moved faster than organisational discipline — which is exactly the argument for the four dashboards in Part Two of this report.

ORGANISATIONAL DESIGN

Discipline, Not Just Ambition

The pattern repeats across technology cycles — Cloud, Agile, Blockchain, RPA — the same trajectory of ambitious launch, resourcing, then momentum loss to **organisational structure failures, not technology limits**. AI is following the same arc, at higher stakes.

01 Enterprise-Level Thinking

Shared data platforms, shared governance, reusable components — not independent departmental programmes that sacrifice compounding value.

02 Locked Execution Cadence

Fixed 90-day delivery cycles, defined deliverables, clear ownership, and governance authority to enforce constraints against scope creep.

03 Culture as Structural Design

Decision rights, rewarded behaviours, and metrics changed through operating structure — not communications and messaging alone.

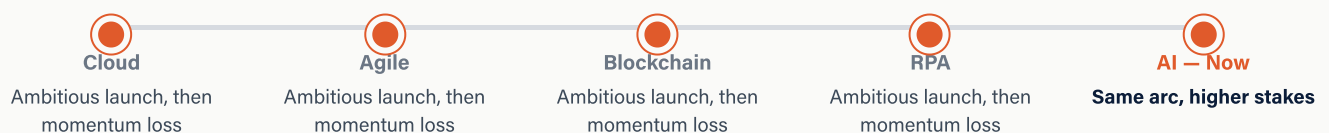
WESTPAC CORE PROGRAMME

Initially cited as a failure; succeeded once restructured around unified strategy, 90-day cadence, and board-level locked scope.

FIRST ABU DHABI BANK MERGER

Framed integration as architecture rationalisation, not systems consolidation — strict constraints delivered the merger on timeline.

THE RECURRING PATTERN



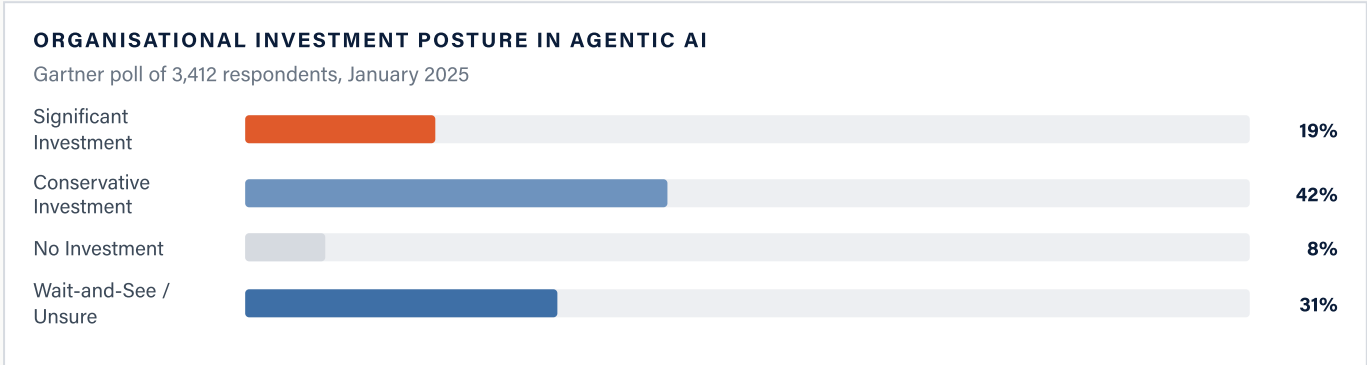
SO WHAT

Convert pilots into platforms. Move from activity to capability through shared infrastructure, locked governance, defined accountability, and a delivery cadence the organisation can actually sustain.

EVIDENCE DEEP DIVE

Why AI Projects Get Cancelled

Gartner's June 2025 prediction that over 40% of agentic AI projects will be cancelled by the end of 2027 was based on a January 2025 poll of 3,412 webinar attendees on their organisation's investment posture. The distribution below is the real starting point for most enterprise agentic AI decisions today.



Gartner names three specific reasons for cancellation: **escalating costs, unclear business value, and inadequate risk controls** — the same three gaps this report's Executive Recommendations address directly. A separate 2025 prediction found that organisations will abandon **60% of AI projects unsupported by AI-ready data** through 2026, and an April 2026 update found **1 in 5 AI projects in IT infrastructure and operations collapses entirely** once in flight.

The market-level consequence is already visible: **42% of companies scrapped most of their AI initiatives in 2025**, up sharply from 17% the year before — a signal that the "wait-and-see" 31% above is the fastest-growing camp, not the shrinking one.

SO WHAT

If your organisation sits in the "conservative investment" 42%, the data says that's the safest place to be — provided the conservatism comes with a locked cadence and a named owner, not just a smaller budget line.

PART TWO • GOVERNANCE

Governing AI at Scale

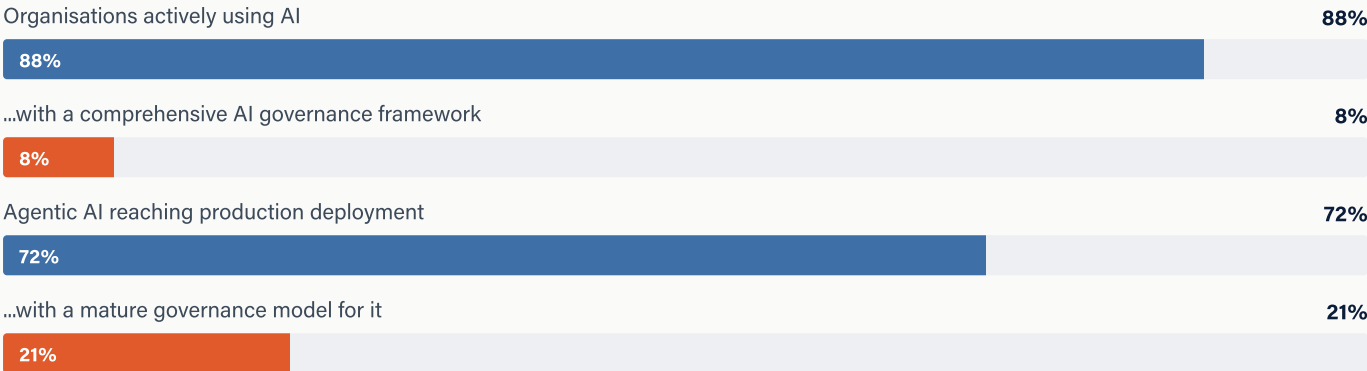
Four dashboards — each paired with its own purpose-and-deployment page — plus a global governance-gap page and a regulatory timeline, for the metrics and mechanisms leaders need once AI becomes core infrastructure.

02

EVIDENCE DEEP DIVE

The Governance Gap

Agentic AI is moving from pilot to production faster than any prior enterprise software category — and governance maturity is not keeping pace. This is the single chart that justifies the four dashboards on the following pages.



40%

of enterprise applications will embed task-specific AI agents by end of 2026, up from under 5% in 2025.

Gartner, 2026

<25%

of executives (of 919 surveyed) have clear visibility into which AI agents are communicating with one another.

Cloud Security Alliance, 2026

36% / 33%

of organisations have adopted ISO/IEC 42001 or the NIST AI Risk Management Framework, respectively.

Stanford HAI, 2026

2%

of small firms have a comprehensive AI governance framework, versus 8% of organisations overall.

Industry governance survey, 2026

SO WHAT

Adoption is not the bottleneck. Governance maturity is — and it is lagging by roughly 50 to 80 percentage points across every metric above. The four dashboards that follow are a starting operating system for closing that gap.

WHY THIS DASHBOARD EXISTS



Enterprise AI Value & Adoption — Purpose & Deployment

This dashboard translates AI spend into the financial language a board already trusts — total cost of ownership, revenue influence, cost savings, and adoption — so AI investment decisions get evaluated with the same rigour as any other capital allocation.

Owner: CFO / Finance Business Partner

Refresh: Monthly (TCO, Adoption) · Quarterly (ROI)

Users: CFO, CAIO, Business Unit Leads

Purpose. It answers the one question every AI budget review eventually asks: is this actually worth what we're spending? It does so by putting TCO, AI-influenced revenue, cost savings, and per-business-unit adoption on a single page, rather than scattered across procurement, HR, and IT reporting.

Why it matters. McKinsey's 2026 survey found only 39% of organisations can attribute any EBIT impact to AI, and fewer than one in five track KPIs for GenAI solutions at all (p.4). Without this dashboard, "AI is working" is an opinion, not a measurement — and opinions don't survive a budget review.

HOW TO DEPLOY IT

- 1

Data sources. Finance/procurement (licensing, implementation cost), HRIS (headcount, hourly cost), and usage logs from each AI tool (licensed vs. active seats).
- 2

Owner. The CFO's office or a Finance Business Partner for AI — this is a financial governance instrument, not an IT report.
- 3

Rollout. Weeks 1–2: build the TCO breakdown from existing procurement data — fastest to stand up. Weeks 3–6: instrument adoption by business unit. Month 2+: layer in revenue and savings attribution once a baseline exists.
- 4

Cadence. Refresh adoption and TCO monthly; refresh ROI and value-per-employee quarterly, once enough data has accumulated to be meaningful.

ESCALATION TRIGGER

Any business unit still below 40% adoption after 90 days should trigger a change-management review — not a tooling review. The technology is rarely the reason adoption stalls at that stage.

C-SUITE MEASUREMENT

Enterprise AI Value & Adoption

Total cost of ownership, AI-influenced revenue, cost savings, and adoption — the live dashboard. Purpose, ownership, and rollout steps are on the preceding page.

Employees: 100

Avg Hourly Cost: \$55

Annual AI Budget: \$180K

Org Adoption Rate: 52%

NET 3-YEAR AI ROI

-86%

conservative model

AI-INFLUENCED REVENUE

\$43K

3-year cumulative

COST SAVINGS

\$59K

3-year cumulative

ORG ADOPTION RATE

52%

active / licensed users

TOTAL 3-YEAR TCO

\$738K

all-in cost

VALUE PER EMPLOYEE

-\$2,121

net annual benefit / headcount

TCO BREAKDOWN (YEAR 1)

Where the \$180K annual AI budget actually goes

\$180K

YEAR 1

Licensing

Implementation

Infrastructure

Governance

Maintenance

\$76K

\$65K

\$25K

\$9K

\$5K

ADOPTION RATE BY BUSINESS UNIT

Target: 70%+ for full ROI realisation (dashed marker)

IT

Sales

Marketing

Operations

Finance

HR

Legal

48%

43%

37%

33%

27%

20%

12%

SO WHAT

A system that looks funded and "live" can still post negative ROI. Track adoption **by business unit**, not just enterprise-wide, to find where value is leaking. Live tool: terencekok.com/resources/ai-value-dashboard

Source: "The Four Dashboards Every Chief AI Officer Must Operate," 22 Jun 2026

16 / 32

WHY THIS DASHBOARD EXISTS



Model Performance & Health — Purpose & Deployment

This dashboard applies production-software monitoring discipline — drift, latency, uptime, incidents — to AI systems that, unlike traditional software, can degrade silently while still returning a plausible-looking answer.

Owner: **MLOps / Platform Engineering**

Refresh: **Real-time (Drift, Latency) · Weekly (Incidents)**

Users: **Engineering, CAIO**

Purpose. It gives deployed models the same instrumentation any production service already has — but tuned for the specific way AI systems fail: gradually, quietly, and without throwing a conventional error.

Why it matters. This report's Nine Failure Modes (p.9) and the Replit incident (p.10) trace back to the same root cause: no one was watching model behaviour in real time. A model with a slowly rising drift score gives no error message — it just gets quietly worse until a customer notices first.

HOW TO DEPLOY IT

- 1 Data sources.** Model inference logs, an MLOps/observability platform (provider-native or dedicated), and an incident ticketing system for rollback tracking.
- 2 Owner.** MLOps or Platform Engineering, with a named on-call rotation — this is an operations dashboard, not a strategy one.
- 3 Rollout.** Week 1: instrument latency and uptime, usually available from existing infrastructure monitoring. Weeks 2–4: add drift detection, which requires a baseline distribution from training or validation data. Month 2: formalise the incident and rollback runbook.
- 4 Cadence.** Drift score and latency monitored continuously; the incident and rollback log reviewed weekly with the CAIO.

ESCALATION TRIGGER

A drift score crossing the alert threshold (0.25 PSI in this report's model) should auto-page the on-call engineer immediately — not wait for the weekly review.

MLOPS MONITORING

Model Performance & Health

Drift detection, latency and uptime, and incident rollback tracking — the live dashboard. Purpose, ownership, and rollout steps are on the preceding page.

Time Range: 7 days

Model Type: Large Language Model (LLM)

Drift Threshold (PSI): 0.25

DRIFT SCORE (PSI)

0.125

within threshold

P95 LATENCY

1.9s

rolling 7-day

SYSTEM UPTIME

99.87%

SLA target 99.9%

OPEN INCIDENTS

0

current

AVG MTTR

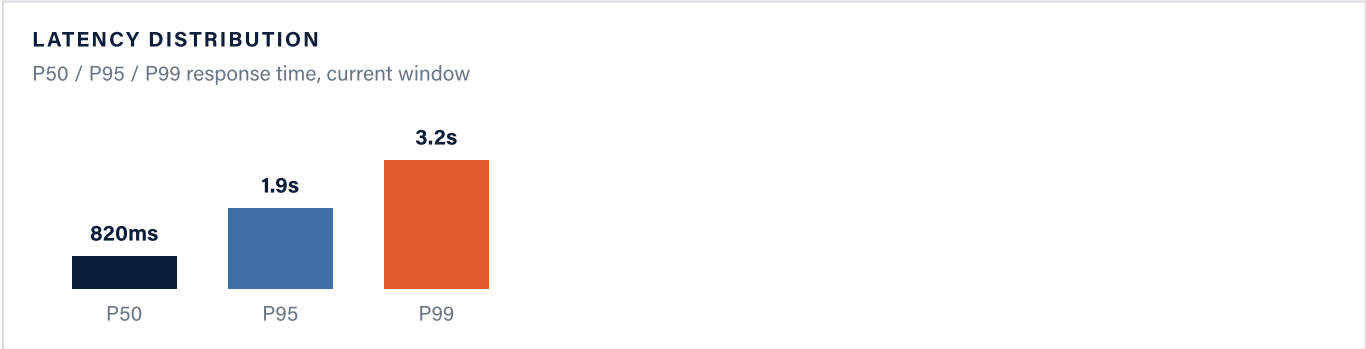
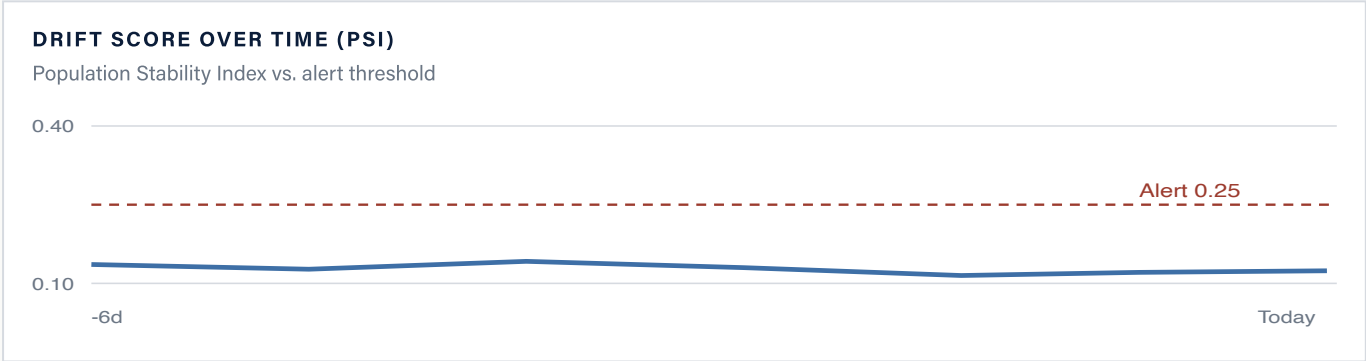
2.8h

mean time to resolution

ROLLBACKS

0

last 7 days



INCIDENT & ROLLBACK TIMELINE

P1

Data drift spike, 4 days ago · MTTR 2.8h · Resolved

P2

Latency SLO breach, 6 days ago · MTTR 1.1h · Rollback · Resolved

SO WHAT

AI models drift quietly, without triggering conventional alerts. Instrument drift and latency with the same rigor as uptime — or you'll find out from a customer first. Live tool: terencekok.com/resources/ai-model-health

WHY THIS DASHBOARD EXISTS



Trust, Risk & Governance — Purpose & Deployment

This dashboard turns three abstract governance domains — bias and fairness, data lineage and privacy, and security — into continuously monitored, audit-ready metrics instead of point-in-time compliance exercises.

Owner: **AI Governance Lead (Legal / Risk)**

Refresh: **Per Release (Bias) · Quarterly (Privacy) · Continuous (Security)**

Users: **Legal, Risk, Compliance, CAIO**

Purpose. It is the difference between "we have a policy document" and "we can prove compliance on demand" — the artefact a regulator, auditor, or board risk committee actually wants to see.

Why it matters. Only 8% of organisations globally have a comprehensive AI governance framework despite 88% actively using AI (p.14). The EU AI Act's Article 50 transparency clock is still running regardless of the Annex III delay (p.24). Governance failures are invisible until an audit or an incident forces the question.

HOW TO DEPLOY IT

- 1

Data sources. Model evaluation results segmented by protected attribute (bias/fairness), a data catalogue or lineage tool (privacy), and a red-team or vulnerability scan mapped to the OWASP LLM Top 10 (security).
- 2

Owner. A named AI governance lead — Legal, Risk, or a dedicated AI Governance function — with technical support from data science for the bias metrics.
- 3

Rollout. Start with whichever regulatory framework already applies to your jurisdiction (this report uses Singapore's IMDA AIGF) and work backward into the three domains from its specific requirements, rather than building generic metrics first.
- 4

Cadence. Bias and fairness testing on every model version change; data lineage and privacy reviewed quarterly or on any new data source; security scanning continuous via automated tooling where possible.

ESCALATION TRIGGER

Any protected group falling below the 80% disparate-impact threshold blocks further deployment of that model version until reviewed — no exceptions for launch-date pressure.

CONTINUOUS ASSURANCE

Trust, Risk & Governance

Bias and fairness, data lineage and privacy, and OWASP LLM security — the live dashboard. Purpose, ownership, and rollout steps are on the preceding page.

Framework: Singapore IMDA AIGF

Risk Tier: High Risk

Maturity: Developing

⚠ Governance Posture: Review Required

IMDA AIGF · High Risk · Developing Maturity

01 · Bias & Fairness Metrics

3/5 GROUPS PASS

Disparate Impact Ratio vs. the 80% Rule (EEOC / NIST AI RMF) — 0.80 threshold



02 · Data Lineage & Privacy

4/5 REQUIREMENTS MET

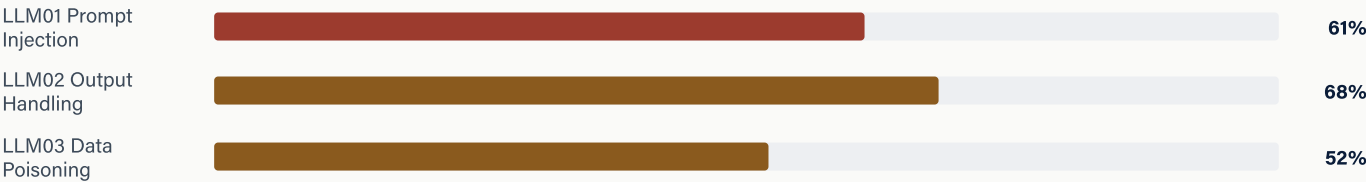
Benchmarked against GDPR Article 5 and MAS TRM data governance standards



03 · Security & Threat Monitoring

0/10 RISKS COVERED

Mapped against the OWASP LLM Top 10 (2023) attack vectors



% detected, current deployment posture

SO WHAT

Governance failures are invisible until an audit or an incident. Treat bias thresholds, data lineage, and OWASP LLM risks as continuously monitored, not annually reviewed. Live tool: terencekok.com/resources/ai-trust-governance

WHY THIS DASHBOARD EXISTS



Human-AI Interaction & Decision Quality — Purpose & Deployment

This dashboard measures whether people and AI are actually collaborating — not just whether AI is deployed — by tracking decision acceptance, automation bias, and augmentation preference against external benchmarks.

Owner: **Business Function Lead + People Analytics**

Refresh: **Continuous (Acceptance) · Quarterly (Trust Survey)**

Users: **Function Leads, HR/OD, CAIO**

Purpose. It catches the failure mode that adoption metrics alone cannot see: people rubber-stamping AI output rather than genuinely evaluating it.

Why it matters. The peer-reviewed research on p.27 shows human-AI collaboration only helps for well-defined, information-rich tasks — and actively hurts complex decision-making unless it is instrumented and reviewed. High adoption without this dashboard can mask exactly the automation bias that erodes decision quality.

HOW TO DEPLOY IT

- 1 Data sources.** Workflow or decision-support tool logs (acceptance/override rate), a periodic pulse survey for AI trust score, and a sample-based quality audit comparing AI-assisted vs. solo-human decisions.
- 2 Owner.** The business function using the AI (e.g. underwriting, customer service), coached by a People Analytics or Organisational Development partner — a workflow-design dashboard, not a purely technical one.
- 3 Rollout.** Instrument the acceptance funnel — AI recommends → meaningfully reviewed → accepted — for a single decision type before expanding. This is the fastest way to catch automation bias early.
- 4 Cadence.** Acceptance and override rates tracked continuously; trust score surveyed quarterly; decision-quality audit sampled monthly.

ESCALATION TRIGGER

An automation bias index above 30% (this report's threshold) triggers mandatory retraining on when and how to override AI recommendations — not a system replacement.

DECISION QUALITY

Human-AI Interaction & Decision Quality

Decision acceptance, automation bias, and augmentation preference — the live dashboard. Purpose, ownership, and rollout steps are on the preceding page.

Industry: Professional Services

Role Level: Management

Decision Type: Complex

✓ Collaborative Maturity: High

Score 77/100 · Professional Services · Management · Complex

DECISION ACCEPTANCE RATE

74%

optimal collaboration zone

AUTOMATION BIAS INDEX

35%

training intervention needed

AI TRUST SCORE

53/100

Edelman baseline 53

AUGMENTATION PREFERENCE

80%

prefer AI as collaborator

DECISION QUALITY LIFT

+16%

vs. solo-human baseline

TIME-TO-DECISION

-22%

faster with AI assistance

AUTOMATION VS. AUGMENTATION SPLIT

Centaur model (human + AI) outperforms either alone by 23% — BCG 2023

75%

AUGMENTED

Human-AI Augmentation

Full Automation

75%

25%

✓ BCG-recommended for complex decisions

ACCEPTANCE RATE BY INDUSTRY

BCG / Stanford HAI 2023 · optimal zone 65–85%

Manufacturing		83%
Healthcare		79%
Prof. Services		74%
Finance		68%
Legal		61%

SO WHAT

High acceptance without review is automation bias, not trust. Watch the acceptance rate **and** the meaningful-review rate together. Live tool: terencekok.com/resources/ai-human-collaboration

Source: "The Four Dashboards Every Chief AI Officer Must Operate," 22 Jun 2026

22 / 32

UNMANAGED RISK

Shadow AI Is Already Inside Your Organisation

90%

of employees already use personal AI tools for work — pasting client details, summarising confidential specs, uploading financial reports, with no processing agreement or audit trail.

**68%**

use unauthorised AI tools at work, up from 41% in 2023.

Second Talent, 2026

66%

used AI despite believing it wasn't permitted under policy.

PagerDuty, 2026

18% / 30%

of organisations have a formal AI security policy / full usage visibility.

Second Talent, 2026

"Ninety percent of your workforce has already decided AI is useful. The question is whether your organisation will shape how they use it, or leave that entirely to them."

RESPONSE 01

Deploy contractually compliant tools with real data-protection controls.

RESPONSE 02

Integrate with existing workflows and systems, not bolt-on portals.

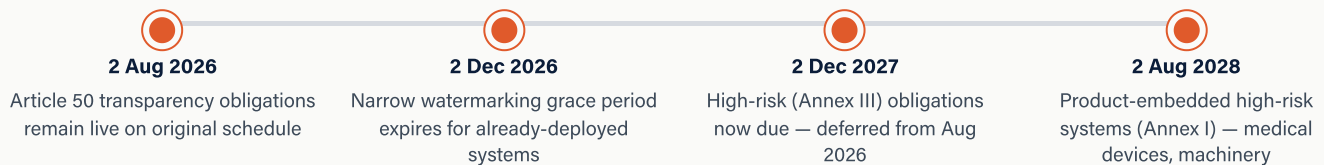
RESPONSE 03

Build visibility through audit trails and usage data, not prohibition.

EVIDENCE DEEP DIVE

The Regulatory Clock

On 29 June 2026, the Council of the European Union gave final approval to the "Digital Omnibus" simplification package — pushing back the compliance deadline for stand-alone high-risk AI systems under the EU AI Act by 16 months. Regulatory delay, however, is not regulatory absence.



Annex III covers AI used in biometric identification, critical infrastructure, education, employment, access to essential services including credit scoring and insurance, law enforcement, migration, and the administration of justice — precisely the categories most enterprise AI governance programmes are built around. The 16-month deferral gives organisations real breathing room on the heaviest compliance lift.

But the **Article 50 transparency requirement** — disclosing to users when they are interacting with an AI system — was not deferred, and remains due on its original 2 August 2026 date. Organisations treating the Omnibus delay as a blanket governance pause are exposed on the one obligation still fully in force.

SO WHAT

Use the 16-month extension to build the governance muscle — Dashboard 03 in this report — rather than to defer the work entirely. The transparency clock is still running, and Singapore's IMDA AIGF, referenced throughout this report's dashboards, has no equivalent grace period.

PART THREE · PEOPLE

The Human Dimension

As machines absorb technical competence, differentiation moves to judgement, curiosity, and accountability—backed by WEF and peer-reviewed collaboration research.

03

DIFFERENTIATION

The Human Advantage in the Age of AI

AI processes information without fatigue, absorbing decades of expertise in seconds. What it structurally cannot do: **feel, choose, or lead**. As technical competence commoditises, LinkedIn's "5 Cs" describe where human value concentrates instead.

 Curiosity Sustained learning	 Courage Decide amid uncertainty	 Creativity Beyond pattern-matching	 Compassion Emotional response	 Communication Human connection
---	--	---	--	---

39% of workers' core skills expected to shift by 2030 (WEF)	83% believe AI will amplify uniquely human skills (Workday)	89% of recruiters rate soft skills as equally important
---	---	---

170M / 92M new roles created vs. roles displaced globally by 2030 — different people, mostly. <i>World Economic Forum, 2025–26</i>	+56% salary premium commanded by workers with specialised AI skills over identical roles without. <i>World Economic Forum, 2025–26</i>
---	---

"In a world where AI can replicate what you know faster than you can acquire it, differentiation emerges from human qualities rather than knowledge accumulation."

EVIDENCE DEEP DIVE

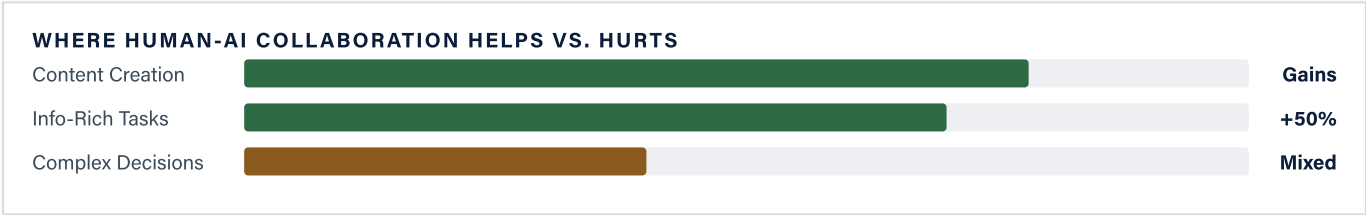
What the Research Actually Shows About Human-AI Teams

A 2026 review of 149 peer-reviewed studies published between 2018 and 2025 gives a more precise — and more useful — answer than "AI augments humans." The honest finding is conditional, not universal.

In well-defined, information-rich tasks, human-AI teams show productivity gains of up to **50%** in field experiments. But the picture reverses for judgement-heavy work: the same body of research finds **performance losses in decision-making tasks**, even as it finds consistent **gains in content-creation tasks**.

The sharpest finding concerns skill mismatch. When an AI system outperforms humans on a task working alone, adding a human to the loop typically makes the combined result **worse** than the AI alone — the human introduces noise, not correction. When humans outperform the AI working alone, the collaboration **beats both** — the AI extends the human's judgement rather than substituting for it. Getting this backwards is the mechanism behind Dashboard 04's "Automation Bias Index."

Across the literature, one variable predicts collaboration quality better than any other: **trust** — specifically, whether the human is willing to act on AI-generated advice, and how much weight they assign it. That willingness depends on interaction design: task procedures, decision authority, and the degree of autonomy granted to the AI all have to match the cognitive style of the human user, not just the technical capability of the model.



SO WHAT

Before mandating AI-in-the-loop for a workflow, ask which category it falls into. For judgement-heavy decisions, the evidence says pair AI with your strongest performers, not your most overloaded ones — the AI amplifies whichever judgement is already in the room.

FUTURE OF WORK

The Entry-Level Opportunity

The traditional entry door has narrowed sharply — but AI tools are so new that **no demographic has accumulated seniority in them yet**, levelling the playing field for those willing to build.

5.6–5.8%

recent-graduate unemployment, among the highest in a decade outside the pandemic spike (NY Fed)

-7%

YoY decline in entry-level hiring specifically, even as overall Class of 2026 hiring rose 5.6% (NACE)

35%

of entry-level roles now explicitly require AI skills

40 → 29

median AI-unicorn founder age, 2021 to 2024 (Antler)

77.2%

of recent grads landed a role within 3 months of graduating — up from 63.3% a year earlier.

ZipRecruiter Annual Grad Report, 2026

81.6% vs 40.7%

hire rate for graduates with work experience during college, versus those without.

ZipRecruiter Annual Grad Report, 2026

THE PARADOX

The traditional entry door is locked, and the market data on this page is not uniformly bad — it's uneven. Building independently is now professionally viable because the playing field genuinely is level: nobody has seniority yet in these tools, and demonstrated experience closes the hiring gap faster than a degree alone. Zach Yadegari built a calorie-tracking app to \$30M in annual revenue at age 19.

CALL TO ACTION

What Leaders Should Do This Quarter

Eight actions that translate this briefing's diagnosis — and its 23 external citations — into an operating agenda.

- 01 Audit AI literacy before funding new infrastructure.** Test actual task performance, not survey confidence.
- 02 Select your next pilot on data availability, workflow fit, and measurability** — not on how strategically ambitious it sounds.
- 03 Weight vendor partnerships over internal builds** for your next pilot — MIT's data shows roughly triple the success rate.
- 04 Stand up at least one of the four CAIO dashboards** before your next board review.
- 05 Replace your shadow-AI ban with a governed, more convenient alternative** — prohibition drives the behaviour underground.
- 06 Lock scope and cadence:** 90-day cycles, named owners, shared enterprise infrastructure over departmental silos.
- 07 Track the EU AI Act's Article 50 transparency deadline (2 Aug 2026)** even though the high-risk obligations were deferred.
- 08 Invest in the 5 Cs** — curiosity, courage, creativity, compassion, communication — as deliberately as you invest in tooling.

FULL REFERENCE LIST

Endnotes & Citations

Every statistic in this report's global-research pages traces to one of the sources below. Original blog citations are listed separately on the following page.

- 1 MIT NANDA / Project NANDA, "The GenAI Divide: State of AI in Business 2025" (2025) — 95% no-P&L-impact finding; vendor-vs-build success rates; methodology of 300+ initiatives, 52 interviews, 153 survey responses.
- 2 RAND Corporation, enterprise AI project research (2025) — 80.3% of enterprise AI projects fail to deliver promised business value.
- 3 Gartner, press release (25 Jun 2025) — >40% of agentic AI projects cancelled by end of 2027; January 2025 poll of 3,412 webinar attendees.
- 4 Gartner, press release (26 Feb 2025) — organisations will abandon 60% of AI projects unsupported by AI-ready data through 2026.
- 5 Gartner (7 Apr 2026 update, via industry coverage) — 1 in 5 AI projects in IT infrastructure and operations collapses entirely.
- 6 Industry analyst data (2026) — 42% of companies scrapped most AI initiatives in 2025, up from 17% in 2024.
- 7 McKinsey, "The State of AI" global survey (2026) — 88% AI usage, 72% GenAI usage, scaling-practice adoption gaps.
- 8 McKinsey, AI Trust Maturity Survey (2026) — governance maturity levels; 39% EBIT-impact finding; 6% "AI high performers."
- 9 Stanford HAI, "The 2026 AI Index Report" — population/organisational adoption rates; \$581.7B investment figure; Transparency Index decline.
- 10 Council of the European Union, "Digital Omnibus" package (approved 29 Jun 2026) — EU AI Act deadline deferrals for Annex I and Annex III systems.
- 11 PagerDuty, workplace AI survey (2026) — 66% used AI despite believing it wasn't permitted; 89% first encountered tools personally.
- 12 Second Talent / SQ Magazine, shadow AI statistics aggregation (2026) — 76% personal-tool usage; 68% unauthorised usage; policy and visibility gaps.
- 13 Gartner (2026) — 40% of enterprise applications embedding task-specific AI agents by end of 2026; 72% agentic production adoption.
- 14 Industry governance survey, via Cloud Security Alliance research notes (2026) — 8% comprehensive governance framework adoption; 74%/21% agentic adoption-vs-governance gap.
- 15 Stanford HAI (2026) — 36% ISO/IEC 42001 adoption; 33% NIST AI RMF adoption.
- 16 Cloud Security Alliance research note (2026) — survey of 919 executives on agent-to-agent visibility.
- 17 World Economic Forum, Future of Jobs research (2025–2026) — 170M/92M job creation-displacement figures; 56% AI-skills salary premium; 22% disruption estimate.
- 18 Federal Reserve Bank of New York; National Association of Colleges & Employers (2026) — graduate unemployment and entry-level hiring figures.
- 19 ZipRecruiter, Annual Grad Report (2026) — time-to-hire and work-experience differential figures.
- 20 Peer-reviewed literature review, human-AI collaboration, 149 studies 2018–2025 (published 2026) — productivity gain ranges; skill-mismatch and trust findings.
- 21 Fortune (7 Oct 2025); OECD.AI Incidents Monitor — Deloitte Australia fabricated-citation report incident.
- 22 Fortune (23 Jul 2025); AI Incident Database, Incident 1152 — Replit AI agent production database deletion.
- 23 Forbes (18 May 2025); Entrepreneur (2025) — Klarna AI customer-service reversal.

REFERENCES

Further Reading

8 Jul 2026	AI STRATEGY
The Avenger Tower Fallacy: Why Your AI Strategy Is Failing Before It Starts	
19 Jun 2026	AI STRATEGY
Why Your Choice of First AI Use Case Will Make or Break the Programme	
28 May 2026	AI STRATEGY
9 AI Deployments Fail (And What to Do About It)	
25 Jun 2026	AI STRATEGY
Why AI Transformation Requires Discipline, Not Just Ambition	
22 Jun 2026	GOVERNANCE
The Four Dashboards Every Chief AI Officer Must Operate	
20 Jun 2026	GOVERNANCE
Shadow AI Is Already Inside Your Organisation	
18 Apr 2026	FUTURE OF WORK
The Human Advantage in the Age of AI	
10 Jul 2026	FUTURE OF WORK
The Entry-Level Job Market Broke. Here's the Opening It Left Behind.	

ABOUT THIS EDITION

This briefing curates and synthesises eight of Terence Kok's 2026 essays into a single executive narrative, layered with 23 external citations from MIT, RAND, Gartner, McKinsey, Stanford HAI, the World Economic Forum, and three documented 2025 case studies. Full articles and interactive tools — including the four dashboards referenced in this report — are available at terencekok.com/blog and terencekok.com/resources.

THE AUTHOR

Terence Kok

Chief AI Officer · Enterprise AI Strategist

25 Years · Asia & The Middle East



Terence Kok is an enterprise AI strategist with twenty-five years leading AI and digital transformation programmes across Asia and the Middle East. He is the builder of **KELIX**, a private cloud RAG platform serving more than 6,000 engineers with zero security incidents, developed through Singapore's IMDA Digital Leaders Programme.

His work centres on three disciplines most AI implementations skip: structured impact assessment before deployment, governance that doesn't slow teams down, and closing the gap between a promising pilot and a system people actually trust in production — the Eight-Dimension framework detailed in his book, *AI at Scale: From Pilot to Production* (2026).

He also created the **Return on Employee (RoE)** framework — used from SME owners to the UN ESCAP AI for Developing Countries Forum — to measure AI value beyond headcount reduction. Terence advises NUS National GRIP, and holds CCTP counter-terrorism certification alongside his AI governance work.

25

years in AI & digital transformation

6,000+

engineers served on KELIX, zero security incidents

3

domains: AI transformation, smart cities, governance

RECENT RECOGNITION

CDO Vision Award 2026 — AIM Media House, for measurable AI & enterprise transformation impact**Google Trailblazer 2.0** — Transformation Award for AI and cloud leadership at scale (2024)**Singapore's Most Influential Data & AI Leaders** — CDO Vision Singapore (2026)**Want to find out where AI actually helps your business?**

Half a day, my methodology, your business — let's see what's actually true about your operations.

terencekok.com

terencekok.com/resources