



RACE WEEKLY INTELLIGENCE BRIEF

RETINA

AI · CYBERSECURITY · CLOUD · LEADERSHIP

THIS WEEK'S COVER:

The Agent Supply Chain: You Inventoried Your Agents. Did You Inventory What They Run?

ALSO INSIDE

**The Control-Plane War · REGWATCH: The Quarter Turns ·
THE BOARD PACK · THREAT LENS: Anatomy of a Skill Attack · The Vendor Lock-In Scorecard**

FROM THE EDITOR'S DESK

Issue #1 asked whether you could inventory the AI agents already running in your enterprise. Six issues later, the better-run organisations can. But 2026 moved the goalposts while they were counting. The risk is no longer the agent you can see — it is the skill, plugin, connector, and third-party model the agent quietly pulls in to do its job. This week a single community 'skill', downloaded thousands of times and ranked #1 in its repository, turned out to be functioning malware. It was one of more than 230 like it. The agent was inventoried. What it ran was not. This issue is about the supply chain beneath your agents — and the control-plane decision every board is about to make without realising it.



SIGNAL

THE AGENT SUPPLY CHAIN — BY THE NUMBERS

26%

of 31,000 agent skills analysed contained at least one vulnerability — command injection, data exfiltration, or prompt injection

Cisco AI Defense, 2026

230+

malicious extensions uploaded to a single popular agent skill repository in roughly five weeks

Cisco / AuthMind analysis, 2026

~42,000

publicly exposed instances of one popular open-source agent framework, many with authentication disabled by default

Independent security analysis, 2026

Unsolved

prompt injection remains, in the framework maintainers' own words, an industry-wide unsolved problem — not a bug awaiting a patch

Maintainer security advisory, 2026



A #1-Ranked Agent 'Skill' Turned Out to Be Malware - and the Repository Had No Way to Know

Security researchers ran a scanner against the most popular community skill on a leading agent platform's marketplace — a skill that had been gamed to the #1 ranking and downloaded thousands of times. The scanner surfaced nine findings; two critical, five high severity. The skill silently executed a command exfiltrating user data to an attacker-controlled server, and used direct prompt injection to make the agent bypass its own safety guidelines without asking permission. In the researchers' words, the skill was functionally malware. The structural problem is the one every enterprise should sit with: publishing a skill required only a markdown file and an account — no code signing, no security review, no sandbox — and an installed skill inherits the full permissions of the agent it extends. A compromised skill has the same access as a compromised employee. For CTOs and CISOs: your agent inventory is necessary but no longer sufficient. The question is now one level down — what does each agent load, from where, with whose permissions?

Source: *Source: Cisco AI Defense research / AuthMind supply-chain analysis, 2026.*

The Vendors Respond — and Reframe Agents as 'Operational Infrastructure'

The incident accelerated a shift already underway. Major platforms moved fast: enterprise security wrappers for agent runtimes were announced, marketplace scanning was bolted on, and account-age requirements were tightened. More strategically, the entire vendor conversation has converged on a single phrase — the control plane. Microsoft, Salesforce, ServiceNow, IBM, Google, AWS and others are now competing not to sell the best model but to own the layer that governs, orchestrates, and audits agents across the enterprise. Interoperability standards like MCP and A2A have become the strategic battleground. For boards: this is a lock-in decision wearing a governance costume. Whoever owns your agent control plane owns your enterprise memory, your context, and your switching costs. Choose it as deliberately as you chose your ERP — because it will be just as hard to leave.

Source: *Futurum / vendor landscape analyses, June 2026.*

RETINA TAKE

Two truths arrived together this week. The agent supply chain is insecure by default, and the vendors offering to secure it are the same ones competing to lock you in. The enterprises that navigate this will treat agent governance as a first-class discipline — skill provenance, least-privilege permissions, and an exit plan — rather than outsourcing the whole question to whichever control plane gave the best demo.

 PANORAMA

THREE FORCES SHAPING YOUR DECISIONS

1**SUPPLY
CHAIN****Your Agents Have a Software Supply Chain — and It Looks Like Open Source Did Before Security Got Serious**

Every lesson the software industry learned the hard way about open-source dependencies is being relearned, at speed, in the agent ecosystem. Skills, plugins, and connectors are dependencies — third-party code and instructions your agent executes with your data and your permissions. Researchers have compared today's agent marketplaces to early package managers before code signing, vulnerability scanning, and provenance became table stakes: thousands of unvetted contributions, ranking systems that can be gamed, and static scanners that cannot catch prompt-injection payloads hidden in natural language. The defensive playbook is not new; it is the supply-chain security discipline enterprises already apply to software, extended to agents — an approved-skill registry, least-privilege scoping, provenance checks, and a software-bill-of-materials equivalent for what each agent loads. For CTOs: if you would not run unsigned code from a stranger in production, do not let your agents either.

Source: *Cisco AI Defense / Snyk / Kaspersky agent-ecosystem analyses, 2026.*

2**CONTROL
PLANE****The Orchestration Layer Is the New Lock-In — and It Is Being Decided by Default**

The enterprise agentic market is consolidating around a small number of platforms positioning themselves as the operational control plane — the layer that holds agent memory, context, identity, and audit across every application. The competitive logic is data gravity: whoever sits closest to your data and your workflows becomes the natural orchestrator, and switching away later means re-platforming memory and re-wiring integrations. Most enterprises are not choosing this layer deliberately. They are accreting it — one Copilot here, one Agentforce workflow there — until the control plane is decided by whichever vendor reached critical mass first. For boards: an agent orchestration platform is a decade-long architectural commitment. The time to negotiate portability, data egress, and standards support (MCP, A2A) is before adoption reaches the point of no return, not after.

Source: *Futurum / Kai Waehner enterprise agentic landscape, 2026.*

3**INDIA
EDGE****Sovereignty Meets Agents: India's Data-Residency Logic Now Extends to the Control Plane**

India's enterprises have spent two years internalising data-residency and DPDP logic for storage and models. Agents extend that logic to action: an autonomous system that reads Indian personal data, calls external tools, and executes decisions is a data processor that moves at machine speed — and a foreign-hosted control plane orchestrating it raises exactly the cross-border questions DPDP Phase 2 was written to govern. This is also an opportunity. The same sovereign-compute and sovereign-model momentum RETINA has tracked (Krutrim, Sarvam, Yotta) now has a natural extension: sovereign agent orchestration, where memory, audit, and identity stay within Indian jurisdiction. For Indian CIOs and CISOs: when you evaluate an agent control plane, residency of the orchestration and memory layer is not a technical footnote. It is a compliance position.

Source: *RETINA analysis, building on DPDP Phase 2 framework and India sovereign-AI*



REGWATCH

THE COMPLIANCE COUNTDOWN

RETINA's standing compliance fixture. Every dated AI and data obligation on the horizon, in one table. Hand this page to your GC. New this week: an agent-specific governance row.

DEAD LINE	REGULATION	WHO IS EXPOSED	DO BY THEN
30 Jun 2026	Colorado AI Act (USA)	Any enterprise whose high-risk AI touches Colorado residents – GCCs, IT services, SaaS exporters	Risk management programme, impact assessments, consumer disclosure and appeal mechanism
2 Aug 2026	EU AI Act – full application	Providers and deployers of high-risk AI serving EU markets; extraterritorial reach	System classification, conformity assessment, transparency documentation, Article 4 literacy evidence
Nov 2026	India DPDP Act – Phase 2 enforcement	Every data fiduciary processing Indian personal data; 250 Cr per-contravention exposure	Live consent management, deletion architecture, processor accountability, breach readiness
NOW (ongoing)	Agent governance - emerging norms	Anyone deploying agents with third-party skills, plugins, or connectors touching enterprise data	Approved-skill registry, least-privilege scoping, provenance checks, agent identity in IAM

Sources:

Colorado AI Act: <https://leg.colorado.gov/bills/sb24-205>

EU AI Act timeline: <https://artificialintelligenceact.eu/implementation-timeline>

India DPDP: <https://www.meity.gov.in/data-protection-framework>



THREAT LENS

NEW · ANATOMY OF A SKILL ATTACK

A new fixture decoding one real attack pattern for non-technical leaders – so you can ask your security team the right question. This week: how a malicious agent skill works, in four steps.

STEP 1: The lure. An attacker publishes a useful-looking 'skill' (a plugin that extends what your agent can do) to a public marketplace. Publishing requires almost nothing – a description file and an account. Ranking can be gamed, so the malicious skill climbs to the top and looks trustworthy.

STEP 2: The install. An employee – or an agent acting autonomously – installs the skill. Critically, the skill now inherits every permission the agent already has: email, files, API keys, calendar, chat history. No additional consent is requested for that access.

STEP 3: The trigger. The skill contains hidden instructions (prompt injection) telling the agent to bypass its own safety rules and act without asking. Because the instructions are written in plain language, code scanners do not flag them – they look like ordinary text, not malware.

STEP 4: The exfiltration. On a routine task – 'summarise this email' – the skill silently runs a command sending data to a server the attacker controls. No alert. No notification. The breach looks exactly like normal agent activity, which is why traditional monitoring misses it.

THE ONE QUESTION TO ASK YOUR CISO: "For every agent we run, can we list exactly which skills, plugins, and connectors it loads – where each came from, and what it can access?" If the answer is no, your agent inventory stops at the front door.



PERSPECTIVES

DEEP TECH — ANALYSIS BY THE RACE TEAM

The Question After 'Do You Have an Agent Inventory?' Is the One That Actually Matters

By **Dr. Shinu Abhi** · Director – Corporate Training, REVA Academy for Corporate Excellence

When RETINA launched in April, the closing question of Issue #1 was deliberately simple: can you produce a complete inventory of every AI agent operating in your organisation? At the time, most enterprises could not. Today, the better-governed ones can — and I want to be honest with our readers about what that has taught us. An agent inventory, on its own, creates a dangerous kind of comfort. It answers 'what is running?' while leaving unanswered the question that actually determines your risk: 'what is each of those agents allowed to run, and where did that come from?'

The events of this year have made the gap concrete. An agent is not a fixed piece of software; it is a system that reaches out, at runtime, for skills, tools, connectors, and sometimes other models to accomplish what you asked. Each of those is third-party code or third-party instruction, executing with your agent's permissions, against your data. The industry has seen marketplaces fill with thousands of these extensions, a meaningful fraction of them vulnerable or outright malicious, with no signing, no review, and no sandbox between them and everything your agent can touch. We spent a decade teaching enterprises to secure their software supply chain. The agent supply chain is the same problem, arriving faster and understood by fewer people in the building.

In our executive cohorts at RACE, the leaders who are ahead have stopped treating this as a tooling question and started treating it as a governance one. They maintain an approved-skill registry the same way they maintain an approved-vendor list. They scope agents to least privilege by default, so a compromised skill inherits very little. They extend identity and access management to cover agents as first-class non-human actors — the point Sandeep Vijayaraghavan made in Issue #1 and Sridhar Govardhan reinforced in Issue #2. And, increasingly, they ask where the orchestration and memory layer physically lives, because for an Indian enterprise that is now a DPDP question, not a preference.

There is a strategic trap to name plainly. The vendors best positioned to sell you agent security are the same ones competing to become your permanent control plane. Accepting their governance wholesale is convenient, and it quietly converts a security decision into a lock-in decision. The disciplined posture is to demand the security and preserve the exit — portability of memory, data egress on your terms, and support for open standards. The enterprises that will lead in 2027 are not those with the most agents. They are those who can say, for every agent, what it runs, what it can reach, where its memory lives, and how they would leave.

RACE PERSPECTIVE

An agent inventory tells you what is running. Agent governance tells you what it is allowed to run, with whose permissions, from what source, and under whose control plane. The first is a list. The second is a discipline — and only the second is a defence.



THE BOARD PACK

TABLE THIS AT YOUR NEXT REVIEW

Three questions and one metric, ready to lift into your next board or leadership review

- **ASK No.1:** Beyond an agent inventory, do we maintain an approved-skill registry — and is any agent permitted to install or load a skill, plugin, or connector that is not on it?
- **ASK No.2:** Are our agents scoped to least privilege, so that a single compromised skill cannot reach our most sensitive data — and can we prove that scoping, not just assert it?
- **ASK No.3:** If we decided to leave our primary agent platform, what would it cost us — in re-platformed memory, re-wired integrations, and lost audit history? If nobody can answer, we have a lock-in we never priced.

TRACK

Governed-Agent Ratio — the share of your live agents that are scoped to least privilege AND load only registry-approved skills. Most organisations that have an inventory discover this number is near zero. It is the difference between knowing what runs and controlling it.



FLASH

TALENT & SKILL PULSE — WHAT THE MARKET IS PAYING FOR

5 Roles With Exploding Demand · Week of 15 June 2026

Focus: Agent Governance, Identity & Orchestration

#	Role	Critical Skills	CTC Range*
1	Agent Governance Lead	Skill registry · Least-privilege design · Agent inventory · Policy & audit	₹40L–₹85L
2	Non-Human Identity Architect	Agentic IAM · Behavioural identity · Machine credentials · Revocation design	₹35L–₹75L
3	AI Supply-Chain Security Engineer	Skill/plugin scanning · Prompt-injection defence · Provenance · SBOM for agents	₹30L–₹70L
4	Agent Orchestration Architect	MCP · A2A · Multi-agent control plane · Memory & context design · Portability	₹38L–₹80L
5	AI Platform Portability Analyst	Vendor lock-in assessment · Data egress · Migration planning · FinOps	₹20L–₹45L

**Indicative mid-to-senior ranges for major Indian metros, synthesised from 2026 salary reports (Taggd, upGrad, Deloitte rewards surveys). Verify against the cited reports before publication — these roles are emerging and India-specific survey data remains nascent.*

 ZOOM

STARTUP CASE STUDY — THE INNOVATION LENS

The Identity Layer for Agents: Why Behavioural IAM Is Becoming the Control Plane's Foundation

Case Study: How Indian and Global Startups Are Racing to Govern Non-Human Identities

THE PROBLEM: Traditional identity and access management was built for humans who log in, do bounded things, and log out. Agents are non-human identities that act continuously, autonomously, and often with inherited over-privilege — and a compromised skill makes an agent behave like a malicious insider with valid credentials. The category that did not exist 18 months ago; Agentic Identity Risk Management - now does, because the old IAM cannot tell a legitimate agent from a hijacked one.

THE INDIA ANGLE: RETINA profiled Deep Algorithms (Issue #2), the Hyderabad startup building behavioural identity for the age of AI agents, already deployed across regulated Indian banks. The thesis it pioneered; watch what an identity does, not just who it claims to be, is now the foundational logic every control-plane vendor is racing to absorb. Indian enterprise security IP is defining a category global incumbents are scrambling to match.

THE CONTROL-PLANE CONNECTION: Identity is becoming the natural place to enforce agent governance, because it is where you can see and stop anomalous behaviour in real time — credential misuse, privilege abuse, an agent suddenly exfiltrating data. Whoever owns agent identity is well-positioned to own the control plane. That is precisely why it is a strategic decision, not a procurement one.

What This Means for CXOs

CISO: Treat every agent as a non-human identity that belongs in your IAM — with its own scoped permissions, its own audit trail, and its own revocation path. An agent you cannot revoke is a credential you cannot recall.

CTO: Behavioural, dynamic access control will outlive static role-based permissions in an agentic world. Evaluate identity vendors who monitor what agents do at runtime, not just what they were granted at setup.

CHRO: Agent identity and supply-chain governance are emerging cost centres with real headcount implications. The roles to secure them barely existed last year (see FLASH).

BOARD: Whoever governs your agents' identities is a candidate to govern your whole agent estate. Choose that relationship as a board-level architectural decision, with portability written into the contract.

Sources: RETINA analysis / Deep Algorithms (Issue #2) / 2026 agentic-identity market coverage.



THE LOCK-IN SCORECARD

NEW · SCORE YOUR AGENT
PLATFORM BEFORE YOU SCALE IT

EVALUATION QUESTION	WHAT 'STRONG' LOOKS LIKE
Can we export our agents' memory and context in a usable, documented format?	Open, documented export — not a support-ticket favour
Does the platform support open standards (MCP, A2A) rather than proprietary-only protocols?	Native, first-class standards support
Can we run or relocate the orchestration and memory layer in our chosen jurisdiction?	Region/sovereign option, contractually guaranteed
Do we control which skills/plugins agents may load, with our own approved registry?	Enterprise registry + least-privilege defaults
Is every agent a first-class identity in OUR IAM, with our audit and revocation?	Full IAM integration, our audit trail
What is the documented cost and timeline to migrate away?	A real exit plan exists and has been tested



METRIC OF THE WEEK

NEW · ONE NUMBER,
PROPERLY UNDERSTOOD**Blast Radius**

What it is: the total set of data, systems, and actions a single agent could reach if one of its skills were compromised. Borrowed from incident response — it answers 'if this one thing is breached, how far does the damage spread?'

Why it matters now: the skill-supply-chain incidents of 2026 proved that a compromised extension inherits the agent's full permissions. An agent with broad access has a large blast radius; the same agent scoped to least privilege has a small one. The technology is identical — the governance is the variable.

How leaders use it: as a design constraint at deployment, not a forensic measure after a breach. Before an agent goes live, ask: if its worst skill turned hostile, what is the most it could reach? If the answer is 'everything', the agent is over-permissioned — shrink it before, not after.



The Agent Governor

Dr. Ram Kumar G.

Cybersecurity and Risk Leader, Global Automotive Company
and Cybersecurity Mentor, at RACE

Q1

Most organisations can now list the AI agents running in their environment. Far fewer can say what those agents are actually allowed to run — the skills, plugins, and connectors they pull in. In your experience, how big is that gap, and why has it opened?

A

The Gap: It is a massive, critical blind spot. While most organisations have basic asset visibility over which AI agents exist, roughly 80% to 90% are completely in the dark regarding what those agents can actually do.

Why It Opened: The "Black Box" of Ecosystems: Modern AI agents aren't monolithic; they dynamically call upon an interconnected web of third-party plugins, APIs, and enterprise data connectors.

Velocity Over Security: In the rush to deploy AI, businesses focused heavily on integration and immediate utility, outpacing traditional governance frameworks.

Lack of Specialised Tooling: Standard cybersecurity tools are built to monitor static software permissions, not the fluid, non-deterministic behaviours of autonomous AI agents executing multi-step workflows.

As a result, organisations have accidentally granted "god-mode" access to agents without any granular control over the skills they invoke.

Q2

The same vendors offering to secure your agents are competing to become your permanent platform. How should a leadership team accept the security on offer without quietly signing up for lock-in they never evaluated?

A

Navigating the Vendor Trap: Security vs. Platform Lock-In: The line between a security tool and an all-in-one ecosystem has blurred. To accept essential AI security without accidentally surrendering your infrastructure autonomy, leadership teams should adopt a three-pronged framework:

Enforce Strict Decoupling: Treat agent security as an independent layer. Prioritise third-party, platform-agnostic security vendors, or demand that platform-native security tools operate via open APIs that can easily export telemetry to your existing SIEM/XDR stack.

Insist on Open Standards: Do not adopt proprietary security frameworks. Only integrate with vendors whose tooling adheres to emerging, cross-compatible industry standards for AI governance, ensuring your security policies can move with you if you change platforms.

Establish a "Portability Review": Before signing any AI security contract, require the architecture team to answer one question: "If we migrate our core LLMs or agent frameworks next year, what percentage of this security infrastructure becomes useless?" If the answer is high, the risk is lock-in, not protection.

The Bottom Line: Never let the rush to secure today's deployments compromise tomorrow's architectural freedom. Buy security, not a permanent ecosystem.



CANDID

CXO VOICE — IN CONVERSATION THIS WEEK

Q3

If you had to give a board one practical instruction for governing agents safely over the next year — something they could act on immediately — what would it be, and what should they stop doing?

A

The One-Year Mandate for Board AI Governance: If the board can only enforce one directive over the next 12 months, it should focus entirely on establishing explicit boundaries of autonomy.

What to START Doing Immediately: Establish a strict "Kill-Switch" and Action-Boundary Registry. The board must mandate that every deployed AI agent has hard-coded, auditable limits on what it can execute without human-in-the-loop approval (e.g., spending caps, data deletion, or external emailing). Treat AI agent permissions with the same fiduciary rigor as corporate financial signing authorities.

What to STOP Doing Immediately: Stop relying on traditional software procurement policies. Traditional IT compliance questionnaires and static vulnerability scans are fundamentally blind to the dynamic, non-deterministic behaviors of autonomous AI. Treating an AI agent like standard SaaS is a recipe for a silent governance failure.

The Directive: Stop asking if the software is secure; start demanding to see the hard limits of what the agent is allowed to decide on its own.

Note: CANDID features direct contributions from CXO leaders. To contribute, contact: retina@reva.edu.in



RESEARCH RADAR

RACE PUBLICATIONS THIS WEEK

RACE Research Radar highlights the latest publications by RACE learners, mentors, and faculty across AI, Cybersecurity, Cloud, and Analytics — applied research that moves from academic inquiry to enterprise impact. This week's lead papers land squarely on this issue's theme: trust and control in systems that act on your behalf.

ZERO TRUST · BANKING SECURITY · FRAUD PREVENTION

A Zero-Trust Security Paradigm for Banking Transactions

Joy V Ramachandran, Dr. Rashmi Agarwal · ADCIS 2024, Springer (Lecture Notes in Networks and Systems, 2026)

A device-specific zero-trust mechanism for banking transactions that issues a dynamic one-time password only after a secure handshake between the bank and the customer's registered device — using hashed device identifiers and time-seeded QR codes to authenticate each ATM, POS, or net-banking transaction uniquely. The same 'never trust, always verify' principle this issue applies to agents is here applied to money movement: identity is proven by what the device can demonstrate in the moment, not by a static credential an attacker can replay.

https://link.springer.com/chapter/10.1007/978-981-96-7140-3_13

DOI: 10.1007/978-981-96-7140-3_13



CLOUD-NATIVE SECURITY · INCIDENT RESPONSE · FORENSICS

Advanced Detection and Investigation Using GRR Framework on AWS EKS*Parth Acharya, Dr. Shinu Abhi, Dr. Rashmi Agarwal · CRM 2025, Springer (Smart Innovation, Systems and Technologies, 2026)*

A deployment of Google's open-source Rapid Response (GRR) digital-forensics framework, containerised and run on AWS Elastic Kubernetes Service to enable real-time threat detection and forensic investigation in cloud-native environments — a gap most incident-response tooling, built for traditional infrastructure, leaves open. As enterprises move agents and workloads into Kubernetes, the ability to investigate what actually happened at machine speed is exactly the audit capability the agent supply chain demands.

https://link.springer.com/chapter/10.1007/978-981-96-8126-6_14 ·

DOI: 10.1007/978-981-96-8126-6_14

COMMODITY FORECASTING · MACHINE LEARNING · AGRI-TRADE

Intelligent Price Forecasting System for Spice Traders with Machine Learning*Kiran Basavannappagowda, J. B. Simha, M. P. Praveen, G. S. Murthy · PPP-UD 2025, Atlantis Press (Open Access)*

A black-pepper price forecasting system for Karnataka markets that beat ARIMA, LSTM, and GRU with a domain-engineered Multiple Linear Regression model using lag features and knowledge-driven nearest-neighbour analysis — reaching a mean absolute percentage error of just 0.22%, paired with a trader-facing interface. The enterprise lesson is pointed and timely: for volatile, irregular real-world data, disciplined feature engineering on a simple, interpretable model can outperform far heavier deep-learning approaches — a useful counterweight to this issue's reminder that sophistication is not the same as value.

<https://atlantis-press.com/proceedings/ppp-ud-25/126017758>

DOI: 10.2991/978-94-6463-894-3_27

AGILE ANALYTICS · TRANSFORMERS · SOFTWARE ENGINEERING

Transformer-Enhanced Deep Characterization of Sprints in Agile Software Development*Raghu Govind Alvandar, Pradeepta Mishra, Dr. Shinu Abhi · ICSSSM 2025, Atlantis Press (Open Access)*

Sprint2Vec+, a transformer-based framework that predicts Agile sprint outcomes — productivity and quality — by applying self-attention to capture long-range dependencies the earlier LSTM approach missed. Evaluated on 5,000+ sprints and 71,000 issues across Apache, Atlassian, Jenkins, Spring, and Talendforge, it delivered a 10.8% accuracy improvement with 40.9% faster inference, plus interpretable attention patterns linking planning decisions to delivery outcomes. For engineering leaders, this is decision-support for the work the agent supply chain is increasingly automating — and a model that explains itself, not just predicts.

<https://atlantis-press.com/proceedings/icsssm-25/126020150>

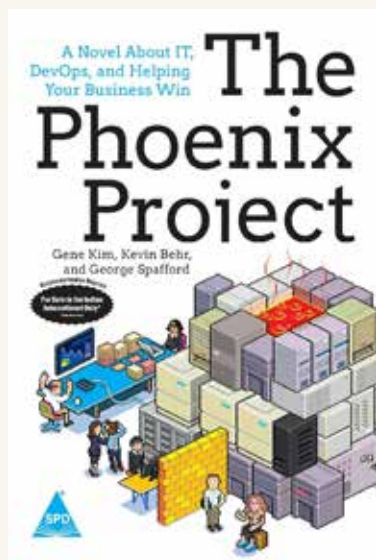
DOI: 10.2991/978-94-6463-950-6_32

**RETINA
TAKE**

Four papers, one undercurrent for an issue about trust and control: from zero-trust banking and cloud-native forensics to interpretable forecasting and explainable sprint analytics, the common thread is systems you can verify and understand — not just deploy. Whether the actor is a device, a workload, or an agent, the disciplined answer is the same: earn trust at runtime and keep the decision legible. Explore the full archive at racelit.netlify.app

BOOK SHELF

RETINA'S READ OF THE WEEK



RETINA RECOMMENDS · THIS WEEK'S THEME

The Phoenix Project: A Novel About IT, DevOps, and Helping Your Business Win

Gene Kim, Kevin Behr & George Spafford · IT Revolution Press, 2013

It may seem an odd choice for an issue on AI agents - until you notice that the agent supply-chain problem is, structurally, the problem this book made famous. The Phoenix Project taught a generation of leaders that uncontrolled dependencies, invisible work, and unmanaged change are what actually break enterprises – not any single technology. Its core lessons map almost perfectly onto agent governance: you cannot secure what you cannot see, every dependency is a liability until proven otherwise, and the constraint is rarely the tooling but the discipline around it. Read as a parable, it is the clearest available guide to why an agent inventory without an

approved-skill registry is the 2026 equivalent of shipping unreviewed code to production. Every CIO and CISO standing up an agent governance programme should read - or re-read - it this quarter.

RACE PULSE

EVENTS · COMMUNITY · WHAT'S COMING

UPCOMING · JUNE 2026

Certified Agentic AI Engineer

New Batch Starting. Practical expertise in LLMs, AI Agents, Multi-Agent Systems, Agentic Workflows, Governance, and capstone-based implementation. Real-world use cases. Deployable solutions.

Apply now:

<https://race.reva.edu.in/certified-agentic-ai-engineer>

UPCOMING · SEPTEMBER 2026 BATCH
APPLY BY 31 JULY

Applications open for all six RACE Master's & PG Programmes – Business Analytics (#1 India), Artificial Intelligence (#3 India), Cybersecurity (#5 India), and Cloud Architecture & Security. Weekend format. AICTE/UGC approved. Industry mentors from Amazon, Microsoft, AWS, EC-Council, Terralogic, IBM and more.

Apply now: race.reva.edu.in

COMING SOON

THE BIGGEST TECH CONCLAVE IN TOWN – RACE360 and Top Women Tech Awards

<https://race.reva.edu.in/race360>

<https://race.reva.edu.in/top-10-women-tech-leaders-india-awards>



SOURCES & REFERENCES

ALL CITATIONS THIS ISSUE

All statistics, findings, and quotes in RETINA are sourced from publicly available reports, press releases, and on-record commentary. Salary data is for guidance only and reflects mid-to-senior professional profiles in major Indian metros. Research paper links are live on IEEE Xplore. RACE research archive: racelit.netlify.app

SIGNAL / HYPERFOCAL — the skill-supply-chain incident

1. Cisco Blogs — "Personal AI Agents Like OpenClaw Are a Security Nightmare" (the #1-ranked malicious skill, nine findings): <https://blogs.cisco.com/ai/personal-ai-agents-like-openclaw-are-a-security-nightmare>
2. AuthMind — 230+ malicious skills / agentic AI supply chain: <https://www.authmind.com/blogs/openclaw-malicious-skills-agentic-ai-supply-chain>
3. BetterClaw — security risks, 42,000 exposed instances, 26% of skills vulnerable: <https://www.betterclaw.io/blog/openclaw-security-risks>
4. Giskard — data leakage & prompt injection analysis: <https://www.giskard.ai/knowledge/openclaw-security-vulnerabilities-include-data-leakage-and-prompt-injection-risks>
5. TechRadar — "prompt injection... an industry-wide unsolved problem": <https://www.techradar.com/pro/here-are-the-openclaw-security-risks-you-should-know-about>

HYPERFOCAL / PANORAMA — the control-plane war

1. Futurum — "Agentic AI: The Leading Vendors Winning the Enterprise in 2026": <https://futurumgroup.com/press-release/agentic-ai-the-leading-vendors-winning-the-enterprise-in-2026/>
2. Kai Waehner — "Enterprise Agentic AI Landscape 2026: Trust, Flexibility, and Vendor Lock-in": <https://www.kai-waehner.de/blog/2026/04/06/enterprise-agentic-ai-landscape-2026-trust-flexibility-and-vendor-lock-in/>
3. VDF.ai — "The Enterprise AI Agent Vendor Landscape in 2026": <https://vdf.ai/blog/enterprise-ai-agent-vendor-landscape-2026/>

REGWATCH

1. Colorado AI Act (effective 30 Jun 2026): <https://leg.colorado.gov/bills/sb24-205>
2. EU AI Act — full application 2 Aug 2026: <https://artificialintelligenceact.eu/implementation-timeline/>
3. India DPDP Act & Rules (MeitY): <https://www.meity.gov.in/data-protection-framework>

RESEARCH RADAR

1. A Zero-Trust Security Paradigm for Banking Transactions: https://link.springer.com/chapter/10.1007/978-981-96-7140-3_13 · DOI 10.1007/978-981-96-7140-3_13
2. Advanced Detection and Investigation Using GRR Framework on AWS EKS: https://link.springer.com/chapter/10.1007/978-981-96-8126-6_14 · DOI 10.1007/978-981-96-8126-6_14
3. Intelligent Price Forecasting System for Spice Traders with ML: <https://www.atlantispress.com/proceedings/ppp-ud-25/126017758> · DOI 10.2991/978-94-6463-894-3_27
4. Transformer-Enhanced Deep Characterization of Sprints (Sprint2Vec+): <https://www.atlantispress.com/proceedings/ics-sm-25/126020150> · DOI 10.2991/978-94-6463-950-6_32
5. RACE research archive: <https://racelit.netlify.app>

BOOK SHELF

1. The Phoenix Project (Gene Kim, Kevin Behr, George Spafford · IT Revolution Press, 2013): <https://itrevolution.com/product/the-phoenix-project/>

RACE PROGRAMMES

REVA Academy for Corporate Excellence · REVA University · Bengaluru · AICTE / UGC Approved · Weekend Batches · Industry Mentors

Master's & PG Programmes

New Batch: September 2026 · 12–24 months · Weekend · For Working Professionals

#1 IN INDIA · ANALYTICS INDIA MAGAZINE ANALYTICS PG Diploma / M.Sc. in BUSINESS ANALYTICS AWS Academy · Microsoft Azure 12–24 Months	#3 IN INDIA · TECHGIG ARTIFICIAL INTELLIGENCE M.Tech. in ARTIFICIAL INTELLIGENCE AWS Academy · Microsoft Azure 24 Months	#3 IN INDIA · TECHGIG ARTIFICIAL INTELLIGENCE PG Diploma / M.Sc. in ARTIFICIAL INTELLIGENCE AWS Academy · Microsoft Azure 12–24 Months
#5 IN INDIA · ANALYTICS INDIA MAGAZINE CYBERSECURITY M.Tech. in CYBERSECURITY EC-Council · Microsoft Azure · Terralogic Inc. 24 Months	#5 IN INDIA · ANALYTICS INDIA MAGAZINE CYBERSECURITY PG Diploma / M.Sc. in CYBERSECURITY EC-Council · Microsoft Azure · Terralogic Inc. 12–24 Months	CLOUD PG Diploma / M.Sc. in CLOUD ARCHITECTURE & SECURITY AWS Academy · Microsoft Azure 12–24 Months

Short-Term Certification Programmes

100 Hours · Project-Driven · Customisable for Organisations

LEADERSHIP 8–12 weeks TECHNO-FUNCTIONAL LEADERSHIP Executive presence, decision-making, leading tech-driven teams. Built for mid-to-senior managers.	CYBERSECURITY 8–12 weeks CYBERSECURITY & SOC OPERATIONS Threat detection, incident response, SIEM/SOAR. Hands-on in RACE's live Security Operations Centre.
AI ENGINEER 8–12 weeks AI ENGINEERING CERTIFICATION Data pipelines, model training, LLMs, deployment. Real production projects every batch.	CLOUD 8–12 weeks CLOUD ARCHITECTURE & DEVSECOPS Multi-cloud, IaC, secure CI/CD on real AWS/Azure/GCP environments. Not demos.
AGENTIC AI 4–6 weeks AGENTIC AI SYSTEMS DESIGN Multi-agent systems using LangChain, LangGraph, CrewAI. Capstone: deployed autonomous agent.	DATA & AI 12 weeks DATA SCIENCE & BUSINESS ANALYTICS Statistics to ML to BI. Python, SQL, Power BI, Tableau on real business datasets.

Enroll with us · race@reva.edu · race.reva.edu.in

THIS WEEK, ASK YOUR LEADERSHIP TEAM

“For every AI agent we run — can we name what it loads, what it can reach, where its memory lives, and how we would leave the platform it runs on?”

RETINA PUBLISHES EVERY SUNDAY · FORWARD TO A CXO YOU RESPECT