

# The Rise of AI-Agents & Agentic Automation for Enterprise Operations

Definitive guide for Market context, strategy, guidelines, governance, and a business-case framework.  
An implementation playbook for Cloud, SRE, IT-Ops and Support teams in the era of enterprise operations

---

Author

**Toshal Khawale**

AWS Community Hero, AWS User Group Leader (Pune)  
Managing Director, PwC.

July 2026



## Current Conditions and the Evolution of Enterprise Operations

Enterprise operations is at an inflection point. Gartner projects that **85% of enterprises will use AI SRE tooling by 2029, up from less than 5% in 2025**. This is one of the fastest technology adoption curves the analyst has forecast in a decade.<sup>1</sup>

The opportunity is real. Mature AIOps deployments cut MTTR by 40 to 58 %, reduce alert volume by up to 95 %, and lift revenue-critical application availability by 15 % in Forrester-commissioned research.<sup>2,3</sup>

The failure risk is equally real. RAND documents an **80% AI project failure rate**. MIT NANDA finds **95% of generative AI pilots deliver no measurable P&L impact**. Gartner reports that only 28% of AI infrastructure projects deliver their expected return, and **57% of I&O leaders have experienced at least one AI initiative failure**. The consistent root cause across all three studies is not technology. It is scope discipline, data readiness, governance, and change management.<sup>4,5,6</sup>

### Three core arguments of this paper are as outlined

**One.** The winning implementation pattern is **hybrid model routing at scale**. Roughly 80 to 90 % of routine ops tasks should run on small language models (SLMs) and open-weight models, hosted self-hosted or in a dedicated tenant, where the token economics and data-residency story hold up. The remaining 10 to 20 %, meaning the genuinely complex reasoning and novel incidents, should route to frontier models via API. Neither pure self-hosted nor pure frontier makes sense. The economics only work when the platform routes intelligently across both.

**Two.** Governance is important and manageable. The frameworks have matured. NIST, ISO/IEC 42001, and OWASP provide clear reference architectures, and vendor implementations are converging. Standard enterprise diligence is sufficient.

**Three.** The economics work at real scale. A base-case business model for a 30-person Cloud, SRE, and Support team at \$150 per hour blended, with a \$30M underlying cloud spend, delivers approximately **\$91M in six-year net value against \$11M in total costs**. Payback lands inside Year 1 and NPV comes in around \$65M at 10 % discount. The full business case is in Section 9.

A fourth argument sits underneath all three: **timing matters**. Agentic operations is a cultural transition, not a tool decision. Organizations that begin now build institutional muscle that late starters cannot easily acquire. Section 1.4 develops this.

# 1. Market Direction. Why AI-Agents, Why Now

Four converging pressures explain the acceleration in AI-agent adoption across IT operations: growing environmental complexity, worsening on-call sustainability, the arrival of production-ready agentic technology, and the strategic cost of waiting.

## 1.1 The Gartner strategic planning assumptions

The January 2026 Gartner Market Guide for AI Site Reliability Engineering Tooling is Gartner's inaugural coverage of this category. It sets a series of adoption assumptions that frame the market:<sup>1</sup>

### GARTNER STRATEGIC PLANNING ASSUMPTIONS (Market Guide for AI SRE Tooling, January 2026)

- By 2029, 85% of enterprises will use AI SRE tooling to optimize operations to meet organizational and customer reliability demands. This will increase from less than 5% in 2025.
- By 2029, 75% of organizations will integrate AI-distilled SRE lessons into product design and delivery, up from 10% in 2025.
- By 2030, 60% of new infrastructure designs will be validated by AI agents that use historical data to identify potential points of failure before development begins.
- By 2029, 90% of organizations will experience an AI-caused outage such as erroneous recommendations, yet continue AI SRE for its speed and scalability gains.

The fourth assumption is worth pausing on. Gartner expects a majority of organizations to experience an AI-caused operational incident yet continue to use AI SRE anyway. This is an unusually candid prediction. It recognizes agentic systems as probabilistic and imperfect, and it directly informs the guidelines in Section 4 and the governance framework in Section 5.<sup>1</sup>

## 1.2 Why now, the four enabling conditions

Gartner identifies three reasons AI-agent tooling has concentrated initially on operational tasks: clearly defined challenges (alert fatigue and manual root-cause analysis are specific, data-heavy problems), return on investment (reduced downtime, fewer engineering hours per incident, prevention of major outages), and availability of training data (operations continuously generate logs, metrics, and traces).<sup>1</sup>

A fourth condition, less discussed but increasingly consequential: **environmental complexity has outrun human expertise**. Modern enterprises run multi-cloud, multi-tenant, multi-tooling estates. Finding engineers who are deep experts across AWS, Azure, GCP, Kubernetes, Terraform, Datadog, Splunk, ServiceNow, and a dozen other tools is difficult and expensive. AI-agents fill this gap by acting as quasi-experts and advisors, giving existing teams the ability to operate at expert level across many tools without hiring more specialists. This is a force-multiplier effect. It is intangible on any single ticket, but it structurally reshapes what a small team can accomplish.

## 1.3 The reliability crunch

Independent research corroborates the case for change. The Catchpoint SRE Report 2025 finds that **~70% of SREs report on-call stress affecting team burnout and attrition**. The 2024 DORA Accelerate State of DevOps Report (39,000 respondents) shows median incident recovery times widening. NeuBird's February 2026 State of Production Reliability survey of 1,039 SRE, DevOps, and IT-Ops professionals documents alert fatigue as a first-class reliability risk. Industry benchmarks put median downtime cost at \$50,000 to \$100,000 per hour.<sup>7,8,9</sup>

Mature AIOps deployments have produced replicable gains: **40 to 58 % MTTR reduction, up to 95 % alert volume reduction, and a 15 % lift in revenue-critical application availability in Forrester-commissioned research**. These are the reference outcomes any implementation should be measured against.<sup>2,3</sup>

## 1.4 The urgency of starting now

Beyond the market direction and the reliability crunch sits a cultural argument that deserves its own consideration.

Agentic operations is not a discrete tool decision. It is a shift in how operations teams work. Engineers who spend the next 12 to 18 months working alongside agents develop instincts for what to delegate, what to review, and how to structure work for machine collaboration. Engineers who wait develop those instincts later, or not at all.

By 2028 or 2029, working alongside AI-agents will be the default operating model for cloud, SRE, and IT-Ops teams. Enterprises that begin the transition now build muscle memory, retained-team competencies, and process changes at a manageable pace. Enterprises that wait compress a two-year cultural transition into six months under competitive pressure, with a talent pool that has already moved to organizations that started earlier.

The strongest reason to start now is not the ROI in Section 9. **It is that the organizations already twelve months in have institutional knowledge and team culture that late starters cannot buy at any price.**

## 2. Understanding AI-Agents for Operations

An AI-agent for operations is not a chatbot with a runbook. It is an autonomous software agent that runs a perceive-plan-act-learn loop. It ingests telemetry and tickets, reasons about the state of a system, executes actions through tools and APIs, and updates its own knowledge based on the outcome.

The 2026 Gartner Market Guide maps the AI SRE tooling category along three capability arcs: incident detection and analysis, remediation and change safety, and reliability engineering. The full capability set, adapted from Gartner's Capability Overview:<sup>1</sup>

| Key use case (business outcome)                                                             | Key capability (technical feature)                                                                                                                                                |
|---------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Accelerate diagnosis and remediation of incidents for SREs, SWEs and IT-Ops teams           | Correlate events to pinpoint root cause from alert noise                                                                                                                          |
| Maintain customer experience in production; protect reliability metrics                     | Analyze telemetry to understand behaviour patterns; reduced MTTR through automated incident response; resource adjustment to maintain customer reliability                        |
| Prevent change risk to defined reliability metrics; protect reliability and user experience | Analyze and continuously refine SLOs and SLIs; proactive deployment risk assessment of changes; autonomous testing and test validation; auto-generation of runbooks and rollbacks |
| Proactively treat reliability as a feature investment in product and platform development   | Prediction of capacity planning by forecasting demand; assessment of architecture resiliency; automated reliability metric discovery                                              |

### 2.1 Three agent classes for enterprise operations

Agentic operations divides into three specialized agent classes that share a common context layer. Each class targets a distinct value pool. Deploying them together on one platform, rather than as siloed products, is what unlocks the compounding value in Section 9.

**AI-SRE.** Incident response and reliability. Alert triage, root-cause analysis, autonomous remediation, change-risk assessment, SLO and SLI tuning, capacity planning.

**AI-Ops.** Operational automation, developer self-service, and closing the gap between visibility and execution. The class most enterprises underestimate. Concrete capabilities include:

- **Faster automation authoring.** Agents write, test, and productionize new operational automations at roughly ten to twenty times human speed, including test coverage. What used to take a senior engineer two days can be delivered in an hour.

- **Centralization of ops execution.** A single agent control plane executes actions across clouds, tools, and repositories. Engineers stop context-switching across five consoles.
- **Proactive advisors.** Agents watch continuously for configuration drift, cost anomalies, security posture changes, and capacity trends. They surface the meaningful signals and either alert or act autonomously depending on the environment.
- **Long-tail task automation.** Backup verification, certificate rotation, environment refreshes, configuration drift cleanup, dependency updates, cross-account log routing. Individually small tasks that collectively consume enormous engineer capacity.
- **CI/CD failure remediation.** Autonomous diagnosis of deployment failures and flaky tests, with proposed fixes generated as pull requests against the affected repository.
- **Developer self-service.** Direct agent-answered queries that deflect roughly 25 % of L1 ticket volume and give developers instant answers without a support handoff.

**AI-Cost Ops.** Continuous cloud FinOps. Waste identification and right-sizing, policy-driven cleanup of idle and orphaned resources, savings-plan optimization, cross-account and cross-region cost attribution.

## 2.2 From visibility to execution

Traditional AIOps stops at visibility. Alert correlation surfaces the right incident. Anomaly detection points at the right dashboard. A cost dashboard tells you spend is high. A scanner tells you a container has a CVE. What happens next is manual, and this is where most operational value is lost.

**Agentic AI-Ops closes this gap by executing on insights.** The pattern is semi-autonomous action, driven by human-in-the-loop where required and full autonomy where safe. Concrete forms:

- **Infrastructure-as-code integration.** Agents draft Terraform, Pulumi, or CloudFormation changes that address the identified issue and open a pull request against the customer's repository. A human reviews and approves the merge. The same pattern applies to Kubernetes manifests, Helm charts, and Ansible playbooks.
- **CI/CD pipeline integration.** Agents propose CI/CD pipeline changes for deployment failures, flaky tests, or performance regressions. The change goes through the normal PR flow with the human as the final gate.
- **ChatOps and ITSM execution.** Agents post remediation candidates in Slack or Teams with a one-click approval that executes through the customer's ITSM or infrastructure APIs. The audit trail is preserved end to end.
- **Continuous drift detection with fixes attached.** Rather than a report saying 'you have 47 configuration drift issues,' the agent posts 47 fixes, each with a one-click apply and a rollback link.

**This execution capacity is what expands the value curve of AIOps into something transformative.** A team that sees an issue and gets a fix ready for review has multiplied its execution capacity many times over. This is the shift that separates modern agentic operations from the previous generation of AIOps tooling.

## 2.3 A representative use-case portfolio

| Agent class | Use case                                       | Business impact                            |
|-------------|------------------------------------------------|--------------------------------------------|
| AI-SRE      | Incident RCA and remediation drafting          | 40 to 58 % MTTR reduction                  |
| AI-SRE      | Alert triage and noise reduction               | Up to 95 % alert volume reduction          |
| AI-SRE      | Change risk assessment and rollback generation | Reduced deployment-caused incidents        |
| AI-Ops      | IaC pull-request generation (semi-autonomous)  | Compress fix cycles from days to hours     |
| AI-Ops      | Long-tail task automation                      | \$2M+ per year at typical enterprise scale |
| AI-Ops      | Developer self-service query deflection        | ~25 % L1 query deflection                  |
| AI-Ops      | Proactive drift and posture monitoring         | Prevent incidents before they occur        |
| AI-Cost Ops | Continuous cloud waste capture                 | 20 to 40 % of underlying cloud spend       |
| AI-Cost Ops | Right-sizing and savings-plan optimization     | 40 to 72 % on committed usage              |

### 3. Best Implementation Strategy

The central recommendation of this paper: **start narrow, run a hybrid model routing architecture, and demonstrate measurable ROI within six to twelve weeks**. Every guideline in Section 4 is designed to protect that trajectory.

#### 3.1 Deployment topology and model routing

Two decisions matter more than most others: where the agent platform runs, and which models it uses for which tasks. These decisions interact.

**On deployment:** self-hosted or BYOC (bring-your-own-cloud) is a strategic requirement for regulated industries, sovereign workloads, and any environment where operational telemetry contains regulated data (PII, PHI, financial transaction records, controlled unclassified information). Model provenance, data residency, and audit trails are enforceable in self-hosted deployments in ways SaaS cannot match. The Cloud Security Alliance's 2026 agentic governance work, Singapore IMDA's Model AI Governance Framework for Agentic AI (January 2026), and NIST's CAISI AI Agent Standards Initiative (February 2026) all treat deployment topology as a first-class governance control.<sup>10,11,12</sup>

**On models:** the winning pattern is hybrid routing. Frontier models offer genuine reasoning depth that SLMs and open-weight models do not yet match, but they are expensive, they add network latency, and most frontier providers cannot be self-hosted. The correct architecture routes routine, high-volume tasks to smaller models running self-hosted or in a dedicated tenant, and reserves frontier model calls for the small fraction of tasks that need advanced reasoning. This is not a compromise. It is what makes the economics work at enterprise scale.

#### 3.2 Model routing architecture

A concrete routing pattern that works in practice:

| Tier                       | Model class                                                            | Share of volume | Rationale                                                                                                                  |
|----------------------------|------------------------------------------------------------------------|-----------------|----------------------------------------------------------------------------------------------------------------------------|
| Tier 1: Routine ops        | SLMs (Phi, Llama-3-8B, Mistral-7B) self-hosted or in VPC               | 60-70%          | Alert triage, log queries, ticket classification, status summaries. Latency-sensitive, data-sensitive, volume-heavy.       |
| Tier 2: Standard reasoning | Open-weight mid-size (Llama-3-70B, Mistral, Qwen) on dedicated hosting | 20-30%          | Standard incident RCA, remediation drafting, cost-optimisation recommendations. Balances capability and cost.              |
| Tier 3: Advanced reasoning | Frontier models (Claude, GPT, Gemini) via API                          | 10-20%          | Novel incidents, complex change-risk assessment, ambiguous multi-step planning, cross-service RCA. Low volume, high value. |

The routing itself has to be platform-owned. A neutral third-party router cannot know that a particular alert is really about a novel dependency failure, or that a particular remediation candidate needs deeper validation. Routing decisions are context-informed, and the context lives in the platform. This is not plumbing. It is a first-class capability.

#### 3.3 Proof of potential. A maturity model

The path from PoC to production runs through three maturity stages. Each stage has a defined expansion gate. No organization should advance to the next stage without meeting the gate.

| Stage | What is happening | Autonomy level | Expansion gate |
|-------|-------------------|----------------|----------------|
|-------|-------------------|----------------|----------------|

|              |                                                                                  |                                        |                                                           |
|--------------|----------------------------------------------------------------------------------|----------------------------------------|-----------------------------------------------------------|
| <b>Crawl</b> | Single use case (e.g. incident RCA on 3 services); human approves every action.  | Human-in-the-loop for all writes       | 50%+ accuracy on evaluation harness; SLA neutral          |
| <b>Walk</b>  | Multi use case across incident, cost, automation; agents draft, humans approve.  | Semi-autonomous in lower environments  | 80%+ accuracy; zero SLA regression                        |
| <b>Run</b>   | Full fleet under governance; autonomous action in production with kill switches. | Autonomous with retained oversight pod | 90%+ accuracy; audit trail in production; rollback tested |

The mistake most organizations make is skipping the Crawl stage or setting no expansion gates at all. Section 4 lists ten concrete guidelines that prevent this.

## 4. Ten Guidelines to avoid PoC Failure

Failure statistics for enterprise AI are consistent across three independent studies. RAND documents an **80.3% AI project failure rate** with 33.8 % abandoned before production, 28.4 % deployed but delivering no value, and 18.1 % deployed but never recouping their cost. MIT Project NANDA finds 95 % of generative AI pilots deliver no measurable P&L impact. Gartner's April 2026 I&O research reports that only 28 % of AI infrastructure projects deliver expected returns and that **84% of failures trace to leadership decisions, not technical limitations**. The following ten guidelines are drawn from the consistent root causes.<sup>4,5,6</sup>

### 1. Define narrow, diverse scope.

Narrow enough to prove value in weeks. Diverse enough to test the platform's generalization. One incident-response use case, one cloud-cost use case, and one automation use case beats fifteen ticket types from the same category. RAND finds vague success criteria present in 73 % of failed projects.

### 2. Choose multi-functional platforms, not single agents.

Every siloed agent creates integration debt, duplicate identity management, and fragmented knowledge. A platform hosting AI-SRE, AI-Ops, and AI-Cost Ops on a shared knowledge layer avoids agent proliferation and unifies governance and audit.

### 3. Insist on an enterprise context layer.

Knowledge graph, memory, and tenant awareness are what separate agents from chatbots. Roughly 60 to 70 % of operational judgment should flow through these layers, not the LLM directly. Gartner's I&O failure research names weak data foundations as the leading structural cause of AI stall.

### 4. Start semi-autonomous with human-in-the-loop.

Autonomous action carries blast radius. Start with AI-drafts and human-approves. Expand autonomy only after the evaluation harness clears 80 % for a task class. Gartner explicitly assumes 90 % of enterprises will experience an AI-caused outage. The question is whether the blast radius is contained.

### 5. Route tasks to the right model. Platform-owned.

Route by task class, latency budget, cost budget, accuracy requirement. Support realtime, async, batch, and target-based patterns. The platform must own routing. A neutral third-party router cannot know what the task actually is.

### 6. Run hybrid: SLMs and open-weight for the majority, frontier for the tail.

Roughly 80 to 90 % of routine ops volume should run on SLMs and open-weight models, self-hosted or dedicated. The remaining 10 to 20 %, the genuinely complex reasoning tasks, routes to frontier models via API. This is what makes economics-at-scale hold.

### 7. Budget for FDEs. Agentic isn't plug-and-play.

Field-deployed engineers integrate the platform to your ITSM (ServiceNow, Jira), observability (Datadog, Splunk, New Relic), ChatOps (Slack, Teams), cloud APIs (AWS, Azure, GCP), and CI/CD. Under-scoping FDE capacity is the most common cause of stalled PoCs.

**8. Change management: process, goals, KPIs.**

Redefine on-call rotations, escalation paths, and performance indicators for L1, L2, L3, the retained team, and the platform team. Org design changes, not just tooling. Gartner's Melanie Freeze names "expected too much, too fast" as the root cause of 57 % of I&O AI failures.

**9. Plan for probabilistic iteration.**

Agentic systems are probabilistic. Start at 50 %-+ accuracy in PoC, iterate to 80 %-+ before pilot, target 90 %-+ for production. Build the evaluation harness in Week 1, not Week 12.

**10. Start fast, start small, but start.**

A focused 6-week PoC with one clear metric outperforms a 6-month evaluation cycle. First measurable value milestone in 90 days is the standard, not a stretch. S&P Global reports 8 months as the median prototype-to-production time. The winning organizations compress that in half.

**Field observation.** A common pattern I have observed is teams committing to autonomous action in production before their evaluation harness is stable. The result is predictable. An early incorrect action forces a rollback of the entire program. Semi-autonomous with human approval for six to twelve weeks builds the confidence and the audit trail needed to expand.

**5. Governance, Security & Compliance**

Security and governance for agentic systems is important, and it is manageable. The field has matured rapidly through 2025 and 2026. NIST published its AI Agent Standards Initiative in February 2026. OWASP published its Top 10 for Agentic Applications in December 2025. ISO/IEC 42001 provides the management-system framework. Enterprise vendors have converged on a multi-pronged approach that combines architectural controls (identity, blast-radius, audit trails) with operational guardrails (approval flows, kill switches, evaluation harnesses gating model updates).

The practical implication for a buyer is straightforward. Standard enterprise diligence is sufficient. The frameworks now exist to make governance a checklist rather than an open-ended exploration. Forrester noted AI-agent risk as a top 2026 CISO concern<sup>13</sup>

, but the same period saw the standards ecosystem move fast enough to give buyers clear reference architectures. The Governance stack in Section 5.1 lists the five domains any credible platform should cover. Section 10's Vendor Scorecard operationalizes verification.

**5.1 The governance stack**

Enterprise-grade agentic deployments cover five governance domains. The framework aligns with NIST AI RMF 1.0, ISO/IEC 42001:2023, the EU AI Act, Singapore IMDA's Model AI Governance Framework for Agentic AI, and OWASP's Top 10 for Agentic Applications. Modern platforms address these through a combination of architecture and guardrails.<sup>14,15,16,11,17</sup>

| Domain                   | Controls (architecture + guardrails)                                                                                                                      | Applicable standard                                     |
|--------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------|
| Identity & authorization | Each agent carries verifiable identity. Least-privilege scopes per agent per task. Time-bounded credentials. MFA-equivalent proof for privileged actions. | NIST CAISI (2026); Singapore IMDA (2026); ISO/IEC 42001 |

|                                     |                                                                                                                                                              |                                                                  |
|-------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------|
| <b>Blast-radius controls</b>        | Per-action approval thresholds. Environment gates (lower-env autonomy vs higher-env HIL). Kill switches. Automated rollback. Rate limits per agent per hour. | NIST AI RMF 1.0 (Manage function); OWASP Top 10 for Agentic Apps |
| <b>Audit &amp; explainability</b>   | Full decision trace for every agent action. Prompt, tool call, and outcome logged. Retention aligned to regulatory horizon. Queryable audit graph.           | EU AI Act Articles 14-15; ISO/IEC 42001                          |
| <b>Data privacy &amp; residency</b> | Self-hosted or BYOC deployment for regulated data. No training on customer telemetry without contract. Regional model routing. PII scrubbing at ingestion.   | GDPR; HIPAA; SOC 2 Type II; EU AI Act                            |
| <b>Model governance</b>             | Model provenance manifest. Evaluation harness gated on updates. Canary rollouts for model changes. Open-weight fallback for regulated workloads.             | NIST AI Agent Standards Initiative; ISO/IEC 42001                |

This stack is not novel or exotic. It is standard enterprise architecture applied to a new asset class. A vendor that cannot articulate how it addresses each of the five domains is not ready for enterprise workloads. A vendor that can is manageable through the same procurement diligence you already apply to any other critical software.

### 5.2 OWASP Top 10 for Agentic Applications

OWASP published its **Top 10 for Agentic Applications** in December 2025. Any platform evaluation should confirm explicit mitigations for the top four categories: prompt injection (both direct and indirect through tool responses), sensitive information disclosure via tool outputs, insecure agent output handling (where downstream systems trust agent output blindly), and excessive agency (agents empowered to take actions beyond the necessary scope). Reputable vendors publish their mitigations and are willing to share red-team results. Ask for them.<sup>17</sup>

### 5.3 Regulated industry considerations

Three regulatory frames deserve specific attention because they shape deployment topology and vendor selection.

|                                                                                                                                                                                                                                                                                                                                                                                                    |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Financial services.</b> For SEC-regulated broker-dealers, agent audit trails must satisfy Books and Records rule 17a-4(f) and FINRA Rule 4511. Self-hosted or BYOC deployment is effectively required to meet WORM retention and access-control obligations. Firms should verify vendors can produce a queryable audit graph on demand, not just a log stream.                                  |
| <b>Healthcare.</b> HIPAA-covered entities and their business associates need explicit BAAs with any LLM provider that processes PHI. Frontier-model API calls that include operational telemetry containing PHI trigger BAA scope. For Tier 3 tasks under the routing architecture, either scrub PHI at ingestion or route Tier 3 through a BAA-covered provider path. Most organizations do both. |
| <b>Public sector and defense.</b> FedRAMP Moderate, IL5, and IL6 workloads require deployment inside authorized enclaves. This effectively rules out SaaS-only agentic platforms. For DoD-adjacent work, self-hosted with an accredited model stack is the only viable pattern. GovCloud or SC2S availability is a threshold question, not a nice-to-have.                                         |

## 6. Common Anti-Patterns

Section 4 lists what to do. Section 6 lists what not to do. These patterns consistently appear in post-mortems of failed AI-agent initiatives.

### ✗ Treating agentic as a chatbot project.

Chatbots answer questions. Agents take actions. When product managers scope agentic like a bot, the integration surface, permissions model, and audit needs get under-specified from day one.

### ✗ Choosing on demos alone.

Every vendor demo works. Insist on a hands-on trial with your telemetry, your ticket taxonomy, and your identity provider. If the vendor cannot show explainable decision traces on your data in the trial, they will not in production.

**✗ No evaluation harness.**

Without a task-specific evaluation harness, "how well is it working" becomes a matter of opinion. Every LLM update, every prompt change, and every agent version needs an automated pass/fail. Build the harness in Week 1.

**✗ No cost monitoring on tokens.**

Token costs at scale can quietly exceed labor savings. Without per-agent, per-task, per-model cost accounting, the ROI narrative collapses at the first CFO review. Model cost should be a separate line item from software cost.

**✗ Boiling the ocean.**

Trying to automate the entire support workload at once guarantees a stalled program. The 80 % of tickets that are AI-addressable can be onboarded in phases. The tail can wait.

**✗ Retained team as shadow support.**

If the oversight pod is doing L1 work in disguise, the platform is not really operating. The retained team supervises the agents, tunes the models, and handles exceptions. It does not do the ticket work.

**✗ Ignoring change management for L1/L2 staff.**

AI-agents change what L1 and L2 engineers do. Career paths, KPIs, and comp models need to move too. Otherwise the people whose cooperation matters most see the platform as a threat.

**✗ Locking into a single model provider.**

A vendor whose architecture assumes one frontier model will pass hidden costs through to you. Pricing changes, API changes, region restrictions all become your problem. Model-agnostic is a strategic property, not a technical detail.

**Field observation.** In multiple engagements the retained oversight team ended up doing L1 work in disguise within three months of go-live. The tell is when the team's weekly reports focus on ticket volume rather than agent performance, false-positive rate, and eval-harness delta. Fix the incentives before you deploy.

## 7. How to Choose an AI-Agent Ops Solution

The market for AI-agent tooling is fragmenting quickly. The following criteria discriminate between platforms that will scale in the enterprise and those that will not.

- **Multi-functional.** AI-SRE, AI-Ops, and AI-Cost Ops on one platform with a shared knowledge layer.
- **Multi-cloud.** No captivity to a single hyperscaler's model stack or observability suite.
- **Model-agnostic.** Support for frontier, open-weight, and SLMs. Model costs separated from software costs. Portable across providers.
- **Context-layer depth.** Knowledge graph, memory, and service graph as first-class objects. 60 to 70 % of decisions flow through the context layer, not the LLM.
- **BYOC or self-hosted.** Required for regulated industries. Also the strongest control for data residency, audit, and model provenance.
- **Extensible and customizable.** SDK, APIs, custom tool integration. The vendor cannot pre-integrate everything you have.

- **Interoperability standards.** MCP or equivalent for tool access. Agent-to-agent protocols emerging under NIST CAISI.
- **Auditability.** Every agent decision has a traceable rationale, queryable and exportable.
- **Ecosystem integrations.** ITSM (ServiceNow, Jira), observability (Datadog, Splunk, New Relic), ChatOps (Slack, Teams), cloud APIs, CI/CD.
- **Governance native.** NIST AI RMF and ISO/IEC 42001 alignment. OWASP mitigations documented.

## 8. The Vendor Landscape

Four architectural approaches are emerging in the AI-agent operations market. Each has structural advantages and constraints that map to different enterprise priorities. The category is young enough that vendor positioning shifts materially each quarter. The summary below reflects publicly available capability data through mid-2026.<sup>18,19</sup>

### 8.1 Observability platform extensions

**Representative vendors:** Datadog Bits AI, Dynatrace Davis AI, Splunk AI, Grafana Assistant.

| Pros                                                                                                                                                  | Cons                                                                                                                                                                                                               | Best suited for                                                                                                                                                                                                                                               |
|-------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Deep integration with existing telemetry. Zero deployment friction for existing customers. Mature UI and workflow. Familiar to on-call teams already. | AIOps-flavored: strong at correlation and postmortem generation, weak at autonomous multi-step action. Bounded to the vendor's own data universe. Vendor lock-in extends from observability to the AI-agent layer. | Organizations that run their entire observability stack on a single vendor (Datadog-only shops, Dynatrace-only shops) and are not yet ready for semi-autonomous or autonomous operations. Fine as a starter tier. Not the answer for full agentic operations. |

### 8.2 Hyperscaler-native agents

**Representative vendors:** AWS DevOps Agent (GA March 2026), Azure SRE Agent, Google Cloud Gemini for SRE.

| Pros                                                                                                                                                                               | Cons                                                                                                                                                                                                                                       | Best suited for                                                                                                                                                                                                                                                              |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Deep cloud-native integration. Aligned to the provider's security and IAM model. Integrated procurement (no separate contract). MCP-based cross-observability integration is real. | Primary optimization is for the home cloud. Multi-cloud reach exists but is secondary. Model choice constrained to the provider's stack. Locking into a hyperscaler agent means locking into the hyperscaler for the life of the platform. | Workloads running only on a single cloud provider with no near-term multi-cloud roadmap. Organizations comfortable with the trade-off: convenience and integration in exchange for permanent vendor lock-in and reduced flexibility on model choice and deployment topology. |

### 8.3 Open-source agent projects

**Representative projects:** K8sGPT, kagent, HolmesGPT, Nudgebee (open-source edition). Nudgebee ships both open-source and enterprise editions.

| Pros                                                                                                                               | Cons                                                                                                                                                                                                                                                        | Best suited for                                                                                                                                                                                                                                                                                           |
|------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Free to try. Full source access. No procurement friction. Strong fit for engineering-led evaluation. Community-driven improvement. | Build and maintain burden is real. The agentic space evolves so fast that internal forks fall behind within a quarter. Enterprise-scale features (multi-tenant governance, audit graph, RBAC, SSO, red-team documentation) are usually missing or immature. | Pilots, PoCs, and internal experimentation, especially at large enterprises with strong platform-engineering benches. Good for teams that want to learn what agentic operations feels like before choosing an enterprise platform. Not recommended as the production destination for regulated workloads. |

### 8.4 Specialist agentic-native platforms

**Representative vendors:** Nudgebee, Resolve.ai, Traversal, NeuBird AI, Cleric, Komodor Klaudia, Sherlocks.ai.

| Pros                                                                                                                                                                                                                                                                                                    | Cons                                                                                                                                                                                                                                                         | Best suited for                                                                                                                                                                                                                                                                          |
|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Truly AI-native architecture, built agent-first from day one, not bolted onto an observability or cloud console. Specialists understand agentic nuances at depth. Vendor-neutral across clouds and observability providers, which minimizes lock-in. Specialization drives faster capability evolution. | Newer market entrants than the observability incumbents. Enterprise track records still building. Require thoughtful integration work into the existing tool stack. Individual vendors vary in coverage across the three agent classes (SRE, Ops, Cost Ops). | Enterprises that want autonomous multi-step action, need multi-cloud and multi-observability neutrality, care about long-term flexibility, and are prepared to invest in a proper implementation. This is the category most aligned with the strategic guidance in Sections 3 through 5. |

8.5 Within the specialist category

Differentiation within the specialist category clusters around three properties: breadth of agent classes, depth of the context layer, and deployment flexibility. The specialists lead in different dimensions. A shortlist typically pairs two or three vendors across these strengths.

Resolve.ai.

Resolve.ai is an autonomous incident resolution specialist.

- **Strength:** depth on autonomous investigation and remediation. Dynamic self-learning knowledge graph targeting 80 % autonomous resolution.
- **Best fit:** teams prioritizing autonomous incident resolution above other agent classes.

Nudgebee.

Nudgebee is a unified multi-agent, AI-agentic Ops platform for enterprise cloud, SRE, and IT-Ops operations. Nudgebee provides three specialist agents on a unified platform. Nudgebee ships in both open-source and enterprise editions.

- **Strenght:** knowledge graph, memory, and service graph as first-class objects for the Context layer.Supports BYOC, self-hosted, and cloud editions deployment.
- **Best fit:** enterprises wanting a unified platform with strong context depth, deployment flexibility, and multi-agent coverage with lowest cost of agentic ops at enterprise scale.

Traversal.

Traversal is a causal-machine-learning specialist for root cause analysis in complex distributed systems.

- **Strength:** causal-ML approach to RCA that reasons about dependencies rather than correlations alone.
- **Best fit:** teams with complex cross-service dependencies needing rigorous causal reasoning.

NeuBird AI.

NeuBird AI is an autonomous SRE agent (Hawkeye) with broad observability-source integrations.

- **Strength:** breadth of observability source integrations. Reads in place from Datadog, Splunk, New Relic, Prometheus, CloudWatch, and Azure Monitor without data migration.
- **Best fit:** teams with heterogeneous observability stacks that want to avoid re-platforming their telemetry.

Komodor Klaudia.

Komodor Klaudia is a Kubernetes-native operations agent.

- **Strength:** deep Kubernetes specialization with cluster-aware reasoning and remediation.
- **Best fit:** Kubernetes-centric operations teams where most workloads are containerized.

Enterprises should shortlist two to three vendors across these strengths for a hands-on trial using the Section 10 scorecard.

**Field observation.** Enterprises that shortlist across all four vendor categories consistently make better long-term decisions than those that shortlist within one category. The hyperscaler agent may win on integration, but comparing it against a specialist during procurement forces the buyer to confront trade-offs they might not otherwise see.

## 9. Business Case Framework

A base-case business model for AI-SRE and AI-Ops deployment at a 30-person Cloud, SRE, and Support team scale, with a \$30M underlying cloud spend footprint. All numbers are anchored to the assumptions table below. An accompanying ROI calculator (available on request) allows re-parameterization for customer-specific values.

### 9.1 Baseline assumptions

| Assumption                            | Value                              | Note                                        |
|---------------------------------------|------------------------------------|---------------------------------------------|
| Support, SRE, and Cloud Ops team size | 30 FTE                             | US enterprise; can scale up or down         |
| Blended engineer rate                 | \$150/hr                           | Fully-loaded US onshore and nearshore blend |
| Annual support labor cost             | \$9.36M                            | 30 × \$150 × 2,080 hrs                      |
| Underlying cloud spend                | \$30M/year                         | Typical for this team size                  |
| Addressable pool (labor + cloud)      | \$39.36M/year                      | Denominator for ROI ratios                  |
| AI coverage ramp Y1 to Y6             | 40 / 50 / 60 / 70 / 80 / 80 %      | Base case adoption                          |
| Effective savings factor              | 0.875                              | 80% full auto + 15% HIL × 50% + 5% human    |
| Cloud waste capture Y1 / Y2+          | 20% / 30%                          | Ramp-up to steady state                     |
| Long-tail ops tasks                   | 500 tasks × 24/yr × (4hr to 30min) | See accompanying ROI calculator             |
| Analysis horizon / discount rate      | 6 years / 10% WACC                 | Standard                                    |

### 9.2 Six-year ROI matrix

| \$M                                            | Y1          | Y2           | Y3           | Y4           | Y5           | Y6           | 6-yr Total    | Beneficiary  |
|------------------------------------------------|-------------|--------------|--------------|--------------|--------------|--------------|---------------|--------------|
| Direct ticket-work savings                     | 3.28        | 4.10         | 4.91         | 5.73         | 6.55         | 6.55         | <b>31.12</b>  | Support cost |
| Cloud cost optimization                        | 6.00        | 9.00         | 9.00         | 9.00         | 9.00         | 9.00         | <b>51.00</b>  | Customer P&L |
| Long-tail automation productivity              | 0.63        | 1.89         | 3.15         | 4.10         | 4.73         | 5.04         | <b>19.54</b>  | Eng capacity |
| Alert triage + noise reduction                 | 0.06        | 0.08         | 0.09         | 0.11         | 0.12         | 0.12         | <b>0.58</b>   | Eng capacity |
| <b>Gross value</b>                             | <b>9.97</b> | <b>15.07</b> | <b>17.15</b> | <b>18.94</b> | <b>20.40</b> | <b>20.71</b> | <b>102.24</b> |              |
| less: Platform license (~\$360K/yr)            | (0.36)      | (0.36)       | (0.36)       | (0.36)       | (0.36)       | (0.36)       | (2.16)        | Platform     |
| less: Oversight team (6 senior FTE @ \$100/hr) | (1.25)      | (1.31)       | (1.38)       | (1.45)       | (1.52)       | (1.59)       | (8.50)        | Retained     |
| less: Year-1 implementation                    | (0.60)      | –            | –            | –            | –            | –            | (0.60)        | One-time     |
| <b>NET VALUE</b>                               | <b>7.76</b> | <b>13.40</b> | <b>15.41</b> | <b>17.13</b> | <b>18.52</b> | <b>18.76</b> | <b>90.98</b>  |              |
| <b>Cumulative</b>                              | <b>7.76</b> | <b>21.16</b> | <b>36.57</b> | <b>53.70</b> | <b>72.22</b> | <b>90.98</b> | <b>–</b>      |              |

**Headline metrics:** Year-1 net value **\$7.76M** | Payback: **inside Year 1** | NPV @ 10%: **~\$65M** | 6-year net value: **\$91M**.

## 10. Illustrated Vendor Evaluation Scorecard

A weighted scorecard operationalizing the criteria in Section 7. Recommended weightings shown. Adjust to your context. Score each vendor 1 (poor) to 5 (excellent). Weighted total ÷ 5 gives a normalized 0-100% fit score.

| Dimension                       | What "excellent" looks like                                              | Weight      | Score |
|---------------------------------|--------------------------------------------------------------------------|-------------|-------|
| Multi-functional coverage       | AI-SRE, AI-Ops, AI-Cost Ops native, shared context layer                 | 15%         |       |
| Multi-cloud and model-agnostic  | AWS, Azure, GCP; frontier, open-weight, and SLM support                  | 10%         |       |
| Context layer depth             | Knowledge graph, memory, service graph as first-class                    | 15%         |       |
| BYOC or self-hosted             | Full BYOC or air-gapped self-hosted; on-prem model option                | 10%         |       |
| Governance native               | NIST RMF and ISO 42001 alignment; OWASP mitigations documented           | 10%         |       |
| Auditability and explainability | Every decision traced, queryable, exportable, retention configurable     | 10%         |       |
| Ecosystem integrations          | ITSM, observability, ChatOps, cloud APIs, CI/CD as first-class           | 10%         |       |
| Extensibility (SDK, APIs)       | Custom tools, custom agents, custom evaluation harnesses                 | 5%          |       |
| Interoperability standards      | MCP or equivalent; agent-to-agent protocols aligned to CAISI             | 5%          |       |
| FDE / implementation support    | Dedicated FDE team, defined implementation methodology                   | 5%          |       |
| <b>Total cost transparency</b>  | <b>Model costs separated from software; per-agent, per-task metering</b> | <b>5%</b>   |       |
| <b>TOTAL</b>                    | <b>Weighted sum ÷ 5 = fit %</b>                                          | <b>100%</b> |       |

**Interpretive guide:** 80 % + fit is a strong candidate for pilot. 60 to 80 % requires customisation or partnership. Below 60 % is high implementation risk.

## 11. Signal vs Noise on Vendor Claims

The AI-agent category attracts a lot of marketing. The following framework distinguishes claims that matter from claims that do not. Use it during vendor pitches and RFP evaluation.

| SIGNAL (what to trust)                                  | NOISE (what to discount)                                          |
|---------------------------------------------------------|-------------------------------------------------------------------|
| <b>Reproducible benchmark on your data</b>              | <b>MTTR reduction %age without denominator or baseline</b>        |
| Documented evaluation harness with pass/fail criteria   | "95 % accuracy" without base rate or evaluation set               |
| <b>Production customer references at your scale</b>     | <b>"10x productivity" claims without task-level breakdown</b>     |
| Audit trail exposed on demand during the demo           | Demo tenant with hand-picked scenarios only                       |
| <b>Published model routing architecture</b>             | <b>Best-case latency without %ile distribution</b>                |
| Public OWASP mitigations documentation                  | Total customer counts without segment breakdown                   |
| <b>Independent red-team results</b>                     | <b>Vendor-designed benchmarks with no third-party validation</b>  |
| Named customers willing to talk on a reference call     | "Trusted by Fortune 500" without specifics                        |
| <b>Per-agent, per-task cost metering in the product</b> | <b>Blended pricing that hides token cost inside software cost</b> |

The rule of thumb: when a vendor answers a question by handing you their evaluation harness and a customer contact, treat it as signal. When they answer with a slide, treat it as noise until you can verify.

## Conclusion

The market direction is clear and Gartner's inaugural Market Guide for AI SRE Tooling projects a 17× adoption jump over four years. Independent Forrester, DORA, and FinOps Foundation data reinforce the case that operations is where agentic AI delivers first-order business impact. The failure literature is equally clear: 80 % of enterprise AI projects fail, and the failures trace to leadership decisions rather than technology limitations.

The path forward is neither cautious retreat nor unconstrained adoption. It is a disciplined implementation program built around narrow-and-diverse scope, a multi-functional platform with a real context layer, hybrid model routing at scale, semi-autonomous operation with human-in-the-loop, aggressive change management, and governance controls aligned to NIST, ISO, and OWASP frameworks. The base-case business model in Section 9 shows that at 30-FTE scale with a typical cloud footprint, a well-executed program returns approximately \$91M over six years against a total investment of \$11M. Payback lands inside the first year and NPV comes in around \$65M at 10 % discount.

The organizations that will lead reliability, cost efficiency, and operational velocity in 2029 are the ones that start focused and iterate. In the words of guideline #10: start fast, start small, but start.

### About the Author

**Toshali Khawale** is an AWS Community Hero and Leader of the AWS User Group in Pune. He professionally serves as a Managing Director at PwC India. His work focuses on enterprise IT operations transformation, Digital Transformation, Cloud Migration, Application Modernisation, cloud economics, and the practical adoption of AI-agentic systems in Cloud, SRE, and IT-Ops functions across regulated industries. Views expressed in this paper are the author's own and do not represent the official position of Employer.

### References

1. Gartner, Market Guide for AI Site Reliability Engineering Tooling, Daniel Betts, Chris Saunderson, Hassan Ennaciri, 26 January 2026 (Gartner ID G00836089).
2. Forrester Research, commissioned Total Economic Impact study on observability and AIOps, 2025. Cited in DevOps.com, "The End of Alert Fatigue," May 2026.
3. INOC, 2026 Event Correlation Guide.
4. RAND Corporation, "The Root Causes of Failure for Artificial Intelligence Projects and How They Can Succeed," 2024.
5. MIT Project NANDA, "State of AI in Business 2025" (published July 2025).
6. Gartner Press Release, "AI Projects in I&O Stall Ahead of Meaningful ROI Returns," Melanie Freeze, 7 April 2026.
7. Catchpoint, "SRE Report 2025."
8. DORA / Google Cloud, "2024 Accelerate State of DevOps Report" and "2025 State of AI-Assisted Software Development Report."
9. NeuBird AI, "2026 State of Production Reliability and AI Adoption Report" (1,039 respondents; February 2026).
10. Cloud Security Alliance, "Agentic AI Governance: NIST AI RMF Agentic Profile," May 2026.
11. Infocomm Media Development Authority (IMDA) Singapore, "Model AI Governance Framework for Agentic AI," January 2026.
12. NIST Center for AI Standards and Innovation (CAISI), "AI Agent Standards Initiative," launched 17 February 2026.
13. Forrester Research, "Top Cybersecurity Threats in 2026," Jitin Shabadu, June 2026.
14. NIST, "Artificial Intelligence Risk Management Framework (AI RMF 1.0)," NIST AI 100-1, January 2023.
15. International Organization for Standardization, "ISO/IEC 42001:2023, Artificial Intelligence Management System," 2023.
16. European Parliament and Council, "Regulation (EU) 2024/1689 (AI Act)," Official Journal of the European Union, July 2024.
17. OWASP GenAI Security Project, "OWASP Top 10 for Agentic Applications for 2026," 9 December 2025.
18. Amazon Web Services, "AWS DevOps Agent, General Availability announcement," 31 March 2026; Microsoft, "Azure SRE Agent, General Availability," 10 March 2026.
19. Flexera, "2026 State of the Cloud Report" (n = 753 cloud decision-makers); FinOps Foundation, "State of FinOps 2026 Report."
20. Google SRE Team, "Site Reliability Engineering: How Google Runs Production Systems" and "The Site Reliability Workbook."

Disclaimer: Gartner does not endorse any vendor, product, or service depicted in its research publications, and does not advise technology users to select only those vendors with the highest ratings or other designation. Gartner research publications consist of the opinions of Gartner's research organization and should not be construed as statements of fact.