



FROM REACTIVE TO AUTONOMOUS

How Agentic AI is reshaping
Enterprise IT Operations

Table of Contents

01. What Makes AI Agentic	4
02. Three Pillars That Determine Success	4
03. What Works and What Fails in Agentic AI Implementation	6
04. Emerging Reality	7
05. About the Author	7

Enterprise AI has moved past the assistant phase.

The next wave is Agentic or Autonomous AI that doesn't wait to be prompted. These systems sense signals, reason through context, and act with minimal human input, making them fundamentally different from the chatbots and co-pilots that preceded them.

For IT operations, the shift is less about capability and more about consequence: what an agent does wrong at 2am, at scale, without a human in the loop, matters in ways that a chatbot response never did..

For IT leaders, agentic AI is not just better automation. It's a move from reactive monitoring to autonomous investigation and resolution, fundamentally reshaping enterprise operations.

Getting there requires more than the right model. It requires clean knowledge foundations, deliberate integration design, and a clear model for where humans stay in control. This piece covers what that looks like in practice: what works, what consistently fails, and where to focus first.



01. What Makes AI Agentic

Traditional automation, a predefined rule determines how a system reacts whenever an event occurs. This means there is a direct association between the event occurring and the system executing a reaction. The entire decision-making process is explicitly programmed beforehand, with every single outcome being predictable.

However, agentic AI works in a different manner. The agents are provided with objectives, not just tasks. They understand their environment through data streams. They analyze what is happening and why. They plan actions and implement them using available tools.

In IT operations, this would translate to agents that not only alert when a threshold is crossed, but also correlate alerts across systems, form hypotheses about the root cause, and recommend courses of action. The implications are very significant: faster detection, lower mean time to resolution, and consistent handling of incidents, no matter who is on call.



02. Three Pillars That Determine Success

Organizations achieving meaningful results from AI adoption share a common approach built on three foundational pillars.



Contextualize: Ground the AI with Knowledge

Agents are only as good as the knowledge they can access. Without understanding and refining existing processes first, organizations automate symptoms, not real problems.

In one implementation, a detection agent monitoring signals and correlating alerts was confidently wrong about severity nearly one-third of the time, flagging SEV2s as SEV0s and missing real SEV0s. The issue wasn't AI; it was biased and

inconsistent historical incident data. The fix wasn't a better model, but clearer severity definitions and retraining on curated examples.

Data quality matters more than model capability. Clean incident history, documented processes, and structured knowledge like Knowledge Graphs enable agents to understand context, not just content. If your data is inconsistent, your AI will be too.



Co-Engineer: Build the Ecosystem

Agentic AI is not an isolated capability; it requires strong integrations, platforms, and feedback loops. The strongest implementations leverage what already exists, design for connectivity (through standards such as MCP and A2A when available), and take a platform-first approach.

Responsible AI requires that AI be designed in from the beginning; security, governance, and

ethics are architectural considerations, not an afterthought.

Most importantly, always design a manual fallback. In critical incidents, orchestration can fail, and teams may need to revert to human-led processes. Integration is often the real bottleneck; fragmented APIs and systems can take longer to align than building the AI itself.



Co-Deliver: Human and AI Together

The strongest model approach positions humans as managers and AI as the workforce; agents handle scale and routine, and humans handle judgment and accountability.

Agents can debug problems by analyzing logs and code changes, but even sound reasoning can lead to errors, so human verification is essential. Confidence measures are helpful, but they don't make decisions. In decision, agents can generate solutions and

rollbacks, but they lack business and organizational context. Key decisions need careful human-in-the-loop checks with complete transparency.

To scale safely, organizations should avoid treating autonomy as a binary. Instead, they must design a graduated path from Level 1 (Human-led/AI-assisted) to Level 5 (Full Autonomy with notification only), allowing trust to scale alongside system capability.

03. What Works and What Fails in Agentic AI Implementation

Experience across implementations reveals consistent patterns. Success is rarely about the model alone, but about how the system is framed, governed, and integrated.

What Works



Start small

Focus on bounded, measurable use cases

Build the foundation

Strengthen knowledge bases before layering AI

Establish feedback loops

Capture human corrections to continuously improve agent performance

Design for failure

Create clear fallback mechanisms for when the agent fails.

One of the most effective ways of thinking about it: consider the agent to be a “junior engineer with perfect memory.” It can remember events, search through logs, and look at runbooks in an instant, but it doesn’t have common sense. The senior engineers are like supervisors, not substitutes. To scale this supervision, use confidence scores as dynamic routing signals for instance, allowing autonomous execution for high-confidence actions, but routing any agent action with a confidence score below 85% to a mandatory human approval gate.

What fails



Automating blindly

Layering AI before understanding the underlying process

Blind trust

Deploying AI without rigorous validation

Speed over safety

Optimizing faster resolution times before building trust in the system

Over-engineering

Using complex agents where simpler programmatic fixes or better alerts would suffice

In one of the triage incidents, the agent picked up the wrong runbook with the same error signature but a different root cause, which resulted in a loss of 20 minutes. The lack of guardrails in speed led to confident misdirection, where clear provenance and confidence thresholds are important. To prevent confident misdirection, agents must provide traceability: a 'reasoning log' that explains exactly which documentation led to the decision, allowing the senior engineer to spot the logic flaw in seconds rather than 20 minutes.

Dashboard metrics may improve, but if engineers still verify everything manually, trust hasn’t been built. In many cases, better alerts, clearer runbooks, and defined escalation paths are enough, no autonomous agent required.

04. Emerging Reality

Organizations adopting agentic AI in IT operations see faster detection, quicker resolutions, consistent incident handling, and durable knowledge captured beyond individual engineers.

The real shift is structural: knowledge is no longer trapped with senior experts. Agents surface past incidents instantly, reduce repeat failures, and enable engineers at all levels to respond effectively.

Transformation, however, takes time. Trust builds gradually, edge cases test systems, and sustained tuning is essential. The question is no longer whether agentic AI will reshape IT operations, but how well organizations adapt to realize its potential through strong knowledge foundations and human-AI collaboration.

“Agentic AI doesn't replace the need for good processes; it makes the cost of having bad processes much more expensive”

05. About the Author



Devanathan Desikan

Devanathan Desikan (Deva) serves as AVP & AI Architect – Digital Services at Movate, where he leads AI-driven offerings for the software delivery lifecycle, delivering impactful AI and engineering interventions. With over 21 years of experience in strategic positions across key IT functions, he has driven capabilities and solutions in Enterprise AI, software and quality engineering, data & analytics, and product management for AI-led platforms and solutions, as well as global technology office initiatives. Deva holds several patents for his innovations in AI.

About Movate

Movate is an Applied AI services company that helps enterprises translate AI ambition into measurable business outcomes. With over 12,000 employees across 20 global locations, augmented by AI collaborator agents, it brings together human expertise and technology to help organizations rethink operating models for efficiency and growth. Powered by Mova iO, its intelligent outcomes platform, Movate delivers practical, scalable, and industry-contextualized solutions that address real business challenges.

For more details, please mail us at info@movate.com