

Securing the Agentic Enterprise

*Threat Modeling, Anomaly Detection, and Governing
Autonomous Multi-Agent Systems*

Khushan Adatiya

Securing the Agentic Enterprise

Threat Modeling, Anomaly Detection, and Governing Autonomous Multi-Agent Systems

Copyright © 2026 Khushan Adatiya

All rights reserved.

First Edition, 2026

No part of this publication may be reproduced, distributed, or transmitted in any form or by any means, including photocopying, recording, or other electronic or mechanical methods, without the prior written permission of the author, except in the case of brief quotations embodied in critical reviews and certain other noncommercial uses permitted by copyright law.

Disclaimer

The information in this book is provided for educational and informational purposes only. It does not constitute professional engineering, legal, regulatory, financial, or security consulting advice. Security architecture decisions depend on facts specific to each organisation, including its regulatory obligations, threat profile, and technical environment, and no book can substitute for assessment of those facts by qualified professionals.

The author and publisher make no representations or warranties, express or implied, regarding the accuracy, completeness, currency, or applicability of the contents. This is a field in which standards, tooling, threat behaviour, and regulatory requirements change rapidly, and material that is accurate at the time of writing may be superseded. Readers should verify current guidance from the relevant standards bodies and regulators before making decisions based on this content.

Neither the author nor the publisher shall be liable for any loss, damage, or disruption arising from the use or misuse of the information contained in this book. Readers are responsible for compliance with all applicable laws, regulations, and organisational policies. Nothing in this book authorises testing against systems for which the reader does not hold explicit written permission.

Views expressed are those of the author alone and do not necessarily reflect the positions of any employer, client, standards body, or organisation mentioned in the text. Research findings, incident reports, and statistics attributed to named organisations remain the work of those organisations and are cited for the purpose of commentary and analysis.

All trademarks, product names, protocol names, and company names appearing in this book are the property of their respective owners and are used for identification and commentary only. Their use does not imply endorsement of this book by the owners, or endorsement of any product by the author.

Printed in the United States of America

Contents

Contents.....	3
Introduction.....	5
The gap this book addresses.....	6
A note on terminology.....	7
Method and evidence.....	8
How the book is organised.....	10
What would falsify the argument.....	12
Part One: The Agentic Attack Surface.....	14
Chapter 1: The Rise of the Agentic Enterprise.....	15
1.1 The Shift from Copilots to Autonomous Agents.....	15
1.2 Quantifying Agent Sprawl.....	19
1.3 Why Legacy Security Tools Fail.....	24
Chapter 2: Threat Modeling the New Execution Boundary.....	29
2.1 Indirect Prompt Injection.....	29
2.2 Memory Injection Vulnerabilities.....	34
2.3 Tool-Calling & Privilege Escalation.....	38
Chapter 3: Agentic Identity: Passports for Non-Human Actors.....	43
3.1 The Machine Identity Crisis.....	43
3.2 Securing the Agentic User Flow.....	47
3.3 AI Agents as the New Insider Threat.....	52
Part Two: Weaponized Autonomy.....	57
Chapter 4: Automated Fraud and Synthetic Profiles.....	58
4.1 Multi-Agent Fraud Infrastructure.....	58
4.2 Scale vs. Sophistication.....	65
Chapter 5: Denial of Wallet (DoW) and FinOps Security.....	71
5.1 The Mechanics of Denial of Wallet (DoW).....	71
5.2 The Resilience Paradox.....	75
5.3 Mitigation: Cloud Cost Arbitrage Layers.....	78
Chapter 6: Secure Agent-to-Agent (A2A) Communications.....	83
6.1 The Risk of Ungoverned Peer-to-Peer Talk.....	83
6.2 Shifting from Shared API Keys to Dynamic A2A Identity.....	89
Part Three: Detecting the Ephemeral.....	94
Chapter 7: Rethinking Anomaly Detection.....	95
7.1 The Failure of User-Centric Baselineing.....	95
7.2 Transition-State Anomaly Detection.....	100
Chapter 8: Monitoring Agent-to-Agent Communication.....	107

8.1 The Agent Communication Visibility Gap.....	107
8.2 Implementing Semantic Firewalls.....	111
Chapter 9: Building the Agentic SOC.....	118
9.1 Ingesting Agent Telemetry.....	118
9.2 Designing SOC Playbooks.....	123
Part Four: Cloud Governance and Trust.....	129
Chapter 10: Zero-Trust Governance Frameworks.....	130
10.1 Sandboxed Tool Execution.....	130
10.2 Continuous Authorization.....	135
Chapter 11: Standardizing Agent Security (NIST & Global Frameworks)..	142
11.1 Navigating the NIST AI Agent Standards.....	142
11.2 Implementing Control Overlays.....	146
Chapter 12: Regulating Autonomous Systems & Compliance.....	152
12.1 The EU AI Act and Agent Autonomy.....	152
12.2 Sharing the Burden.....	156
Part Five: The Consumption-Based Economic Threat.....	162
Chapter 13: Consumption-Based Enterprise Risk: Hallucinations, Cost Arbitrage, and Alternative Pricing.....	163
13.1 Hallucinations as a Financial Threat.....	163
13.2 Consumption-Based Metering Vulnerabilities.....	168
13.3 Enterprise Alternatives & Structural Shifts.....	172
Part Six: Architectural Blueprint for Executive Action.....	176
Chapter 14: Architectural Blueprint for Executive Action and Long-Term Strategy.....	177
14.1 Immediate Technical Directives for CISOs.....	177
Conclusion.....	187
Acknowledgments.....	196
About the Author.....	197
Glossary.....	199
Works Cited.....	202

Introduction

Enterprise software has, for most of its history, been obedient in a narrow and useful sense. A program did what its authors compiled it to do. When it failed, it failed along paths that a competent engineer could trace back to a line of source. Security architecture grew up inside that assumption. Firewalls filtered traffic because traffic had predictable shapes. Static analysers scanned code because the code that would run was the code that had been written. Access control granted permissions to identities because identities were durable things attached, in the last instance, to a person who could be called on the telephone and asked what they had been doing at two in the morning.

That assumption has quietly stopped holding. A large fraction of the software now being deployed inside enterprises writes its own instructions at runtime, decides for itself which systems to contact, retains what it learns across sessions, and completes multi-step objectives without a person watching any individual step. The industry calls these systems agents. The name is unfortunate, because it suggests something small and helpful, a clerk running errands. What has actually been introduced into corporate networks is a class of non-deterministic executor that holds valid credentials, operates continuously at machine speed, and derives its next action from text it read a moment ago, including text written by people who do not work for the company.

This book is about the security consequences of that change. My argument, stated plainly at the outset so that the reader can weigh the evidence against it, has three parts.

First, the security boundary of an agentic system is semantic rather than syntactic. The question that matters is no longer whether a request is well formed and carries a valid credential. Both of those things will be true of nearly every damaging action an agent takes, because the agent is authorised and the request is legitimate in every mechanical respect. The question that matters is what the request means in relation to the task the agent was given. Almost none of the enterprise security stack, as deployed in 2026, is capable of asking that question.

Second, autonomy converts several long-standing and mostly theoretical vulnerabilities into practical, high-frequency ones. The confused deputy problem, described in the computer science literature in the 1980s, was for decades a curiosity that required an unusual arrangement of privileges to exploit. An autonomous agent with tool access is a confused deputy by construction, sitting permanently in production, waiting for instructions from whatever text it happens to retrieve. Something similar has happened to the distinction between instructions and data, a separation that compilers and operating systems have enforced for half a century and that language models dissolve as a condition of their operation.

Third, the economic surface has widened alongside the technical one. When compute was bought in units of servers, an attacker who wanted to hurt an organisation financially had to break something. Under consumption-based billing, an attacker can hurt an organisation financially by making its own software work harder. The recursive, self-correcting behaviour that engineers deliberately design into agents so that they recover from transient errors becomes, under a token meter, a mechanism for converting a single malicious prompt into a five-figure invoice. Security and financial operations are no longer separate concerns with a courtesy relationship. They are the same concern viewed from two directions.

The gap this book addresses

There is no shortage of writing on artificial intelligence and risk. Most of it falls into two categories, and both leave the practitioner without much to do on Monday morning. The first category is model safety: alignment, refusal behaviour, bias, the properties of the model considered as an artefact in isolation. That work is serious and necessary, and it is largely irrelevant to the enterprise security engineer, because the model is the one component in the stack that the enterprise did not build and cannot inspect. The second category is governance at the level of principle: frameworks, maturity models, statements that organisations should maintain oversight and manage risk. Also serious, also necessary, and also difficult to translate into a firewall rule.

The territory between those two bodies of work is where the incidents actually happen. It is the territory of identity issuance for software that has no login page, of telemetry schemas for systems whose most interesting state is a chain of natural-language reasoning, of authorisation decisions that must be made mid-task rather than at session start, of anomaly detection for processes that have no working hours and no fixed address. This book stays in that territory. It treats the model as given, treats the organisation as a real place with a budget and an audit committee, and asks what has to be built.

The intended reader is the person who will be asked to sign off on an agentic deployment: security architects, platform and infrastructure engineers, heads of security operations, chief information security officers, and the compliance and risk professionals who now find themselves needing to understand execution graphs. A reader with a background in distributed systems will find the terrain familiar in structure, since much of what follows is a story about identity, state, and observability under conditions of high concurrency and partial failure. A reader coming from governance will find the technical material explained rather than assumed. Neither reader is asked to write code, though both are asked to understand what the code does.

A note on terminology

Vocabulary in this field is unsettled to a degree that causes real confusion in architecture reviews, so the terms used throughout this book are fixed here. An agent, as the word is used in the following chapters, is a software process that receives an objective expressed in natural language, decomposes that objective into steps without external direction, selects and invokes tools to execute those steps, and evaluates the results before deciding what to do next. Three properties are load-bearing in that definition: the objective is under-specified relative to the actions required, the plan is produced at runtime, and the loop continues without a person approving each iteration. A system missing any of the three may be valuable, and may even be marketed as an agent, but it does not raise the security questions this book is about.

An orchestrator is the software that maintains agent state, holds the tool registry, manages the execution loop, and coordinates multiple agents where more than one is involved. It is the closest thing an agentic deployment has to a control plane, and the argument of Chapters 8 and 10 depends on the orchestrator being treated as security infrastructure rather than as application code. A tool is any capability exposed to an agent through which it can affect a system beyond its own context: a database query interface, an HTTP client, a message-sending function, a code execution sandbox. The set of tools available to an agent is the set of things that agent can do, and the security review of an agentic deployment is largely a review of that set.

A session is a bounded execution context beginning when an agent receives an objective and ending when the objective is completed or abandoned. The session is the unit at which task scope can be declared, authority can be delegated, and behaviour can be baselined, which makes it the natural unit of enforcement. Persistent memory refers to state deliberately retained across sessions, typically in a vector store or structured summary, and it is treated in this book as a separate asset class from session context because its security properties are different in kind, as Chapter 2 explains.

Two terms are used with more caution than is common elsewhere. Autonomy is a matter of degree rather than a binary, and where a claim depends on how much autonomy a system has, the text says which decisions are being made without human review. Alignment appears rarely, because the property that matters operationally is whether an agent's actions stay within the authority it was granted, which is a question about architecture rather than about the disposition of the model.

Method and evidence

Claims in this book are attached to named sources wherever a source exists. Where a figure comes from a specific study or industry benchmark, that study or benchmark is identified in the text so the reader can check it. Where a mechanism is described but no public measurement of its prevalence exists, the text says so rather than manufacturing

a number, which is a discipline worth stating explicitly given how much writing in this field consists of confident statistics with no traceable origin.

Two limitations follow from the subject matter and should be acknowledged early. The first is recency. Agentic security is a field in which the standards bodies published their initial positions during the period this book was being written. Specific control catalogues will move. The structural arguments should outlast them, because they follow from properties of the architecture rather than from any particular vendor's implementation. The second limitation is disclosure asymmetry. Organisations that suffer agentic incidents have strong reasons not to describe them, and the public record is therefore biased toward incidents that were visible from outside, involved a research team willing to publish, or happened to a company large enough that concealment was not an option. The reader should treat published incident counts as a lower bound and nothing more.

A word on what this book deliberately does not do. It does not recommend products. Vendors are named where a named vendor published research, proposed a specification, or documented an architecture that the argument depends on, and naming them is a matter of attribution rather than endorsement. The controls described here are architectural patterns, and every one of them can be implemented with more than one set of components. Readers looking for a shopping list will be disappointed, which is the correct outcome, because the shopping list would be stale before the book reached them.

The reader should also be warned about a specific failure of reasoning that recurs in this field, because this book relies on analogy repeatedly and analogy is where the errors come from. Agentic systems resemble several things that came before. They resemble microservices in their concurrency and their machine identities. They resemble insider threats in their position inside the trust boundary. They resemble unsanctioned software adoption in their pattern of arrival. Each comparison illuminates something and misleads about something else, and the discipline this book tries to maintain is to state both halves. Where an analogy is introduced in the chapters that follow, the paragraph

that introduces it is usually followed by a paragraph explaining where it breaks, and the second paragraph is generally the more useful of the two.

A related caution applies to the numbers. Survey research in security has structural problems that no methodology fully corrects. Respondents are self-selected, definitions of terms such as agent vary across organisations in ways that make counts non-comparable, and the vendors who commission the research have an interest in the answers. The figures reported in this book are the best available at the time of writing, and several of them come from vendor-commissioned surveys, which is stated where it is the case. They are used to establish orders of magnitude and directions of movement, and the arguments are constructed so that they survive substantial error in any individual figure. Where an argument would collapse if a statistic were wrong by a factor of two, that fragility is noted.

Finally, a remark on the relationship between this material and the broader public argument about artificial intelligence. Much commentary treats the risks of autonomous systems as speculative, concerning capabilities that do not yet exist and harms that have not yet occurred. Whatever one concludes about that debate, it is a different debate from this one. The systems discussed in the following chapters are deployed today, in production, at the organisations most readers work for. The incidents described are documented incidents with dates. The failure modes are observable in current architectures using currently available tooling. Nothing in this book depends on any assumption about future capability, and a reader who believes that the longer-term concerns are overstated should still find the operational argument intact, because it rests entirely on properties that present systems already have.

How the book is organised

Part I establishes the attack surface. Chapter 1 describes the architectural transition from human-supervised copilots to autonomous execution graphs, quantifies how far deployment has outrun governance, and explains why the existing defensive stack is structurally blind rather than merely under-configured. Chapter 2 builds the threat

model, working through indirect prompt injection, memory injection into persistent context stores, and privilege escalation through tool-calling. Chapter 3 turns to identity, which is where most of the practical remediation work begins, and sets out the delegation model that makes agent actions attributable.

Part II follows autonomy into the hands of attackers. Chapter 4 examines multi-agent fraud infrastructure and the synthetic identity pipelines that mature profiles over weeks before using them. Chapter 5 treats denial of wallet as a first-class attack class and explains why the resilience properties engineers value make the problem worse. Chapter 6 addresses agent-to-agent communication, where the trust assumptions inherited from service meshes turn out to be badly wrong.

Part III is about detection. Chapter 7 explains why user-centric baselining fails against continuous cloud-native processes and proposes transition-state modelling in its place. Chapter 8 covers monitoring of inter-agent traffic and the semantic firewall pattern. Chapter 9 describes the telemetry a security operations centre needs to ingest and the response playbooks that replace device isolation when there is no device.

Part IV addresses governance. Chapter 10 applies zero-trust principles to tool execution and continuous authorisation. Chapter 11 covers the standardisation work under way at the national level and how agentic controls map onto existing catalogues. Chapter 12 examines regulatory exposure and the way the cloud shared responsibility model divides liability when an agent exceeds its scope.

Part V takes up the economics directly, treating hallucination as a financial event and cost arbitrage as an attack. Part VI closes with a prioritised architectural plan for the executive who has to decide what to do first, and with an assessment of which of these problems are engineering problems and which are structural ones that will not be engineered away.

The chapters are written to be read in order, since the threat model in Part I is assumed by the detection design in Part III and the governance argument in Part IV. A reader with a specific immediate problem can start with the chapter that names it and follow

the internal references backward. The one section I would ask every reader not to skip is Chapter 3, because identity is the substrate on which every other control in this book depends, and an organisation that cannot say which agent did a thing cannot implement any of the rest.

What would falsify the argument

A book that advances a thesis owes the reader an account of what evidence would count against it, and there are three developments that would substantially undermine the argument made here. The first would be a durable architectural separation between instructions and data at the model layer. If a future model architecture enforced a genuine privilege boundary, so that content arriving through a designated channel could not influence behaviour regardless of its phrasing, the injection chapters of this book would become historical. Research in this direction exists, and the dual-model and capability-restriction proposals discussed in Chapter 2 are early moves toward it. None of them currently provides a guarantee, and a guarantee is what would be required.

The second would be a demonstration that semantic enforcement can be performed cheaply and reliably at production volume. Much of the architecture recommended in Part III assumes that evaluating the meaning of an action is expensive enough to require tiering and imperfect enough to require defence in depth. If small, fast classifiers achieve high accuracy on intent evaluation, the control problem simplifies considerably and the elaborate session-scoped machinery described here becomes overengineering.

The third would be empirical evidence that the attack classes described in Part I and Part II remain rare in practice despite widespread deployment. The argument of this book treats the mechanisms as sufficient reason for concern, on the grounds that the economics favour attackers and the detection gap is wide. If several years pass with agentic deployment continuing to grow and incident rates staying flat, the correct conclusion would be that some property of real deployments constrains these attacks in a way the structural analysis fails to capture. Readers encountering this book after such a period should weigh the incident record more heavily than the reasoning.

Setting out these conditions is a matter of intellectual hygiene, and it also has a practical use. Each of the three describes a capability an organisation should watch for, because the arrival of any of them would justify reallocating security investment away from the controls this book recommends. Security architecture is a set of bets on how a threat environment will develop, and the bets should be revisited when the odds change.

One final framing note. It would be easy to read a book of this kind as an argument against deploying autonomous systems. It is not that. The productivity case for agentic automation is real, the deployments are happening at a rate that no security function is going to reverse by memorandum, and a security programme built on the premise that the technology should not exist will be routed around within a quarter. The useful question is narrower and harder: what has to be true about an agentic deployment before a competent security professional can defend it in front of a board. The chapters that follow are an attempt to answer that question with enough specificity to be actionable and enough honesty to admit where the answer is still being worked out.

Part One: The Agentic Attack Surface

Chapter 1: The Rise of the Agentic Enterprise

The transition described in this chapter happened faster than the vocabulary used to describe it. Terms such as copilot, assistant, and agent are used interchangeably in vendor material and in a good deal of internal architecture documentation, which obscures a distinction that carries most of the security weight. This chapter draws that distinction precisely, measures how far the resulting deployments have run ahead of governance, and explains why three categories of established defensive tooling cannot see what these systems are doing.

1.1 The Shift from Copilots to Autonomous Agents

The first generation of enterprise generative systems was stateless and conversational. A person typed a request, a model produced text, and the person decided what to do with it. The model had no memory beyond the current window, no ability to act on anything outside the chat surface, and no capacity to continue working when the person stopped typing. Every consequential action passed through a human hand. Whatever one thought of the output quality, the security properties were straightforward, because the model was an advisory component and the enforcement point was the employee's own judgement.

Modern deployments have inverted that arrangement. The current pattern places the model inside a loop: it receives an objective, decomposes it into steps, selects tools to execute those steps, observes the results, and revises its plan accordingly. Orchestration frameworks such as LangGraph and AutoGen exist to construct and manage exactly this loop, maintaining state across steps and coordinating multiple specialised models against a shared objective. The human supplies the objective and, in a well-designed system, reviews the outcome. The intermediate steps, which are the ones that touch databases and external endpoints, are unobserved by default.

A useful way to hold the difference is to consider where the program counter lives. In conventional software, control flow is a property of the source: a developer wrote the

branches, and the set of reachable states was fixed at compile time. In an agentic system, control flow is generated at runtime from natural-language context. The set of reachable states is not enumerable in advance, because it depends on text the system has not yet read. This is the property that breaks the assumptions underneath most security tooling, and it is worth being clear that it is a property of the design rather than a defect in any particular implementation.

Several research groups have described the resulting arrangement using an operating system analogy, with the model acting as a kernel that arbitrates access to storage, context, external interfaces, and local runtimes. The framing was presented at academic venues including COLM in 2025 and has since become common in architecture discussions. As an analogy it repays attention, because it makes an uncomfortable implication explicit: a kernel is the component that everything else trusts, and this particular kernel makes its scheduling decisions by interpreting text. Every property that makes an operating system kernel a sensitive piece of software applies here, with the added complication that the arbitration logic cannot be audited by reading it.

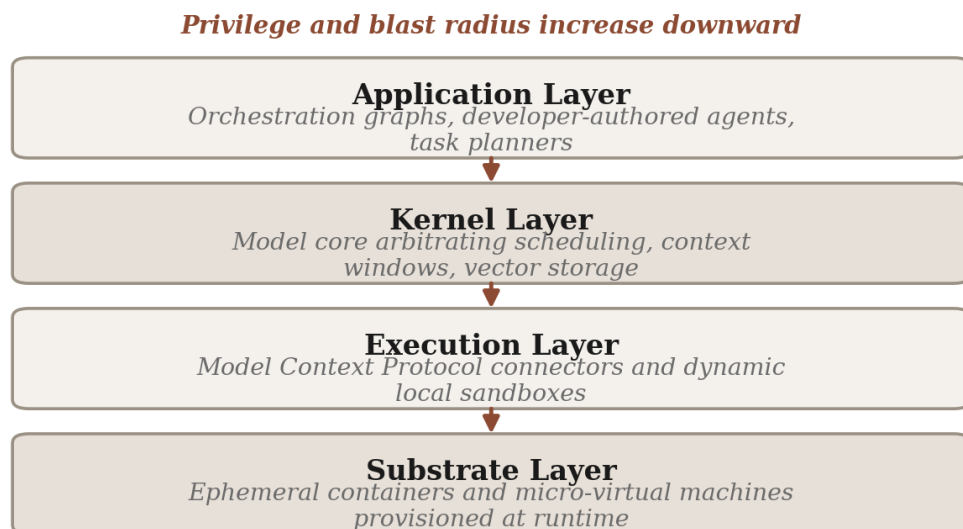


Figure 1.1 The layered structure of an agentic runtime. Each descending layer widens the blast radius of a compromised planning loop.

The execution layer deserves particular attention. To complete tasks that involve running code, transforming data, or interacting with systems that have no clean interface, orchestrators provision short-lived execution environments: containers, sandboxed shells, temporary virtual workspaces. These environments are created in response to a plan the model generated moments earlier and destroyed when the task completes. The security perimeter is therefore not a fixed boundary that can be drawn on a network diagram. It comes into existence at runtime, in a shape determined by a non-deterministic process, and disappears before most scanning tools would have got round to inventorying it.

Connectivity has standardised in parallel. The Model Context Protocol has been widely adopted as a common interface between agents and the tools they call, which has solved a real integration problem and created a real security one. A standard connector layer means that granting an agent access to a new system is a configuration change rather than an engineering project. Anything that lowers the cost of granting access raises the rate at which access is granted, and the protocol as specified carries no native mechanism for constraining what an authenticated caller may do once connected.

System attribute	Human-supervised copilot	Autonomous agentic pipeline
Operational control	Sequential prompt cycles, each initiated by a person	Stateful execution graphs that continue without supervision
Logic construction	Deterministic code paths fixed before deployment	Plans synthesised at runtime from natural-language context
Integration boundary	Human validation at every system interface	Automated tool-calling through standardised connectors
Execution environment	Static application servers with known inventory	Ephemeral containers provisioned in response to a plan
Identity model	Credentials belonging to the operating employee	Non-human workload identities acting on delegated authority

Table 1.1 Structural differences between supervised and autonomous deployment patterns.

Read across that table and a pattern emerges that should be uncomfortable for anyone who has spent a career building controls. Every row describes the removal of a place

where a control used to sit. The prompt cycle was a checkpoint. The fixed code path was an object of review. The integration boundary was where a human noticed that something looked wrong. The static server was where the agent-based scanner lived. The employee credential was what made the audit log meaningful. None of these were formal security controls in most organisations. They were incidental properties of the architecture that happened to constrain the damage a compromised component could do. Autonomy has removed them together, and nothing has been put in their place by default.

It is worth being concrete about how the loop behaves in failure, because the failure modes are where the security properties become visible. A conventional program encountering an unexpected condition either handles it along a path the developer wrote or terminates. Those are the only options, and both are legible afterwards. An agent encountering an unexpected condition reasons about it. If a query returns an error, the agent may conclude that the schema has changed and attempt to discover the new one. If discovery fails, it may reformulate the problem, select a different tool, or escalate to a larger model. Each of these decisions is locally sensible and none of them was written down in advance. The system does not fail so much as improvise, and the improvisation happens in production, at speed, with whatever credentials the agent holds.

This improvisation is the property engineers value most and the one security teams should examine most closely. It is what allows an agent to complete a task despite a malformed input or an intermittent dependency, which is a genuine operational advantage over brittle automation. It is also the mechanism by which a small perturbation becomes a large deviation. An agent that has decided the schema changed will pursue that hypothesis, and the actions it takes in pursuit of a wrong hypothesis are as authorised as the actions it would have taken in pursuit of a right one. Chapter 5 shows how the same behaviour becomes a financial attack surface under consumption pricing, and Chapter 7 shows why it defeats anomaly detection built on rates and volumes.

One further architectural detail is worth recording before moving on, because it determines where telemetry can be collected. In most orchestration frameworks the reasoning trace, meaning the sequence of intermediate conclusions the model produced while deciding what to do, exists in memory during execution and is discarded afterwards unless the deployment explicitly captures it. The trace is the only artefact that explains why an action was taken. A deployment that logs tool calls without logging the reasoning that produced them retains the what and loses the why, which is sufficient for billing reconciliation and insufficient for incident investigation. This is a configuration decision usually made by a platform engineer in the first week of a project, typically without consulting anyone in security, and it determines whether an investigation eight months later is possible at all.

Multi-agent arrangements add a further layer that single-agent descriptions omit. Where an objective is complex, orchestrators commonly decompose it across specialised agents: one that plans, several that execute parts of the plan, sometimes one that reviews the output of the others. Coordination happens through messages passed between them, and the messages carry both instructions and results. The security consequence, treated properly in Chapter 6, is that the boundary between components inside the system has the same properties as the boundary at its edge. A message from one agent to another is text arriving in a context window, and text arriving in a context window is the mechanism by which everything in Chapter 2 happens. Architects accustomed to service meshes tend to treat inter-component traffic as trusted by default, which is a reasonable position for deterministic services and an unsafe one here.

1.2 Quantifying Agent Sprawl

The deployment curve has been steep enough that the usual governance mechanisms have not had time to engage. According to the Gravitee State of AI Agent Security 2026 Report, corporate agent fleets roughly doubled in early 2026, with close to 38 percent of surveyed organisations managing more than a hundred active agents, against an average of 37 in late 2025. The absolute numbers matter less than the shape of the curve, which

is the shape of a technology being adopted by business units directly rather than being rolled out by a platform team on a schedule.

The governance figure from the same body of research is the one that should command attention. Only 14.4 percent of organisations report that their production agents went live with the full approval of the security or information technology function. The remaining 85.6 percent were provisioned by line-of-business teams pursuing an immediate productivity gain, without a security review, a threat model, or in many cases an entry in any inventory that the security team can query.

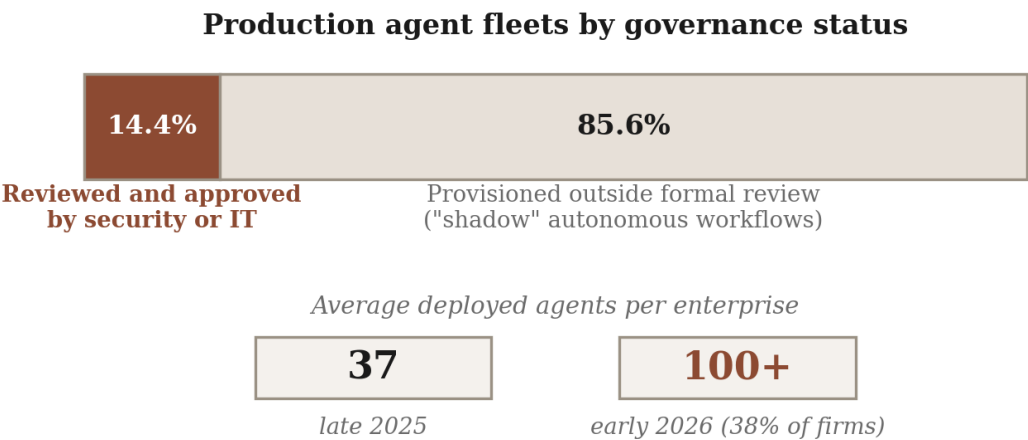


Figure 1.2 Governance status of production agent fleets, with the growth in fleet size over the same period.

Practitioners will recognise the pattern from an earlier decade. Unsolicited software-as-a-service adoption followed the same route: a team with a problem, a credit card, and a vendor promising a solution by Friday. The analogy is instructive but incomplete, and the difference between the two cases is where the risk lives. An unsolicited project management tool sits inert until someone opens it. An unsolicited agent reads files, calls endpoints, writes to systems, and repeats the cycle on a schedule, at machine speed, with no requirement that anyone be watching. The idle-versus-active distinction is the entire security story. A tool that waits has a small attack surface most of the time. A tool that works continuously has its full attack surface exposed continuously.

The permissions granted to these workflows compound the problem. Business units provisioning their own automation are optimising for the thing working, not for least privilege, and the fastest way to make an agent work is to give it broad access and stop debugging. Field reporting on shadow agent deployments describes agents granted direct access to source repositories, internal messaging archives, and production database systems, typically through credentials issued for a human user or through a service account whose scope was never narrowed after the initial proof of concept.

Governance metric	Late 2025 baseline	Early 2026 status
Average deployed agents per enterprise	37	More than 100 in 38 percent of organisations
Formal security or IT approval	Under 10 percent of deployments	14.4 percent of production agents
Mean detection delay	241 days for a standard breach	247 days for a shadow AI compromise
Incremental breach cost	Baseline	An additional 670,000 US dollars on average

Table 1.2 Governance and cost metrics for unmanaged agent deployment.

The cost figures translate the governance gap into a number a finance committee will understand. Industry analysis of breach economics puts the incremental cost of a shadow AI compromise at roughly 670,000 US dollars above the standard baseline, bringing the average total to about 4.63 million. The premium has a single dominant driver: a detection delay of approximately 247 days, which is what happens when the compromised component was never inventoried, produces no telemetry that any collector is configured to receive, and performs its damaging actions using credentials that every downstream system regards as valid.

Two hundred and forty-seven days is worth sitting with. It is eight months during which an agent with database access, a poisoned instruction in its persistent memory, and a valid token has been operating inside the network while every dashboard reported green. The reason it reported green is that from the perspective of the monitoring stack, nothing anomalous occurred. An authorised service made authorised calls to systems it

was permitted to reach. That is precisely what the detection problem consists of, and it is the subject of Part III.

There is a further complication that the survey data does not capture well. Agents can provision other agents. An orchestrator handed a complex objective may spin up specialised sub-agents to handle parts of it, each with its own execution context and its own tool access. An inventory taken on Monday can therefore be wrong by Wednesday, not because a human deployed something new but because a system did. Any inventory approach built on periodic scanning inherits this problem. Inventory has to be a property of the runtime, emitted by the orchestrator as agents are created, rather than a report generated by walking the estate and counting.

There is a temptation, on reading figures like these, to conclude that the answer is prohibition: a policy stating that agents may be deployed only with security approval, enforced by whatever means the organisation has available. Policies of this kind have been written in most large enterprises during the past eighteen months. The evidence that they work is thin, for a reason that has nothing to do with the quality of the policy. The productivity gain from an agentic workflow accrues to the business unit that deploys it, is visible within days, and is measured in hours of work removed. The risk accrues to the organisation, is invisible until an incident, and is measured in a probability. That asymmetry has defeated every prohibition-based approach to technology adoption in the history of corporate computing, and there is no reason to expect a different outcome here.

The workable alternative is to make the approved path cheaper than the unapproved one. If a business unit can obtain an agent with scoped credentials, a logging pipeline, and a tool registry by filing a request that is answered the same day, the incentive to build one privately mostly evaporates. If the approved path takes six weeks and produces something less capable, teams will continue to route around it and will become better at concealing the fact. This is a platform engineering problem more than a policy problem, and organisations that have made progress on shadow deployment have

generally done so by building the paved road rather than by strengthening the prohibition. The controls described in Chapters 3, 9, and 10 are considerably easier to implement on a platform that agents are actually deployed through.

A second observation about the numbers concerns what they measure. Counting agents is a harder exercise than it sounds, and different organisations answering the same survey question are frequently counting different things. Some count orchestrator deployments, some count distinct configured personas within one orchestrator, some count running instances, and some count integrations. An organisation reporting four agents may be running four hundred processes. This matters for defence because permission grants, credential rotation, and audit scope attach to the running instance rather than to the configuration, and an inventory built on configurations will systematically understate exposure. The practical recommendation, developed in Chapter 9, is to derive inventory from runtime telemetry rather than from a register that people maintain by hand.

The procurement pattern behind these deployments explains a good deal about their security posture, and it is worth describing because it suggests where intervention is cheapest. Agentic capability arrives in an organisation through at least four doors, and only one of them passes through anything resembling a technology review. The first is deliberate platform adoption, in which a team stands up an orchestration framework and builds agents against it. This is the visible case, the one security teams tend to have in mind, and it is a minority of the total. The second is feature activation inside software the organisation already owns, where a vendor adds agentic capability to an existing product and enables it by default. No procurement event occurs, no contract is renegotiated, and the security review that covered the original product covered a product with different properties.

The third door is developer tooling. Coding assistants that execute commands, run tests, and modify repositories are agents by the definition used in this book, and they typically hold credentials to source control, build systems, and sometimes production. They

arrive through individual developer choice, often on personal accounts, and their access derives from the developer's own entitlements. The fourth is browser extensions and desktop utilities that act on a user's behalf inside authenticated sessions. These hold no credentials of their own, which is why they escape credential-based inventory entirely, and they can perform any action the logged-in user can perform.

Ranking these by risk-to-visibility ratio produces an uncomfortable ordering. The deliberate platform deployment is the highest-capability case and the most visible, which means it is the one most likely to receive controls. The vendor feature activation and the developer tooling are moderately capable and nearly invisible. The browser-resident case is capable of anything the user is and appears in no inventory whatsoever. An organisation that has secured its orchestration platform and stopped has addressed the door it could see. Discovery work aimed at the other three, which is mostly a matter of examining what employees have installed and what vendors have enabled, tends to produce more risk reduction per hour spent than any amount of hardening applied to the platform that was already under management.

1.3 Why Legacy Security Tools Fail

It is tempting to treat the visibility problem as a configuration failure, something that would be resolved if security teams tuned their existing tooling more carefully. That reading is comfortable and wrong. The three main categories of enterprise defensive tooling fail against agentic systems for structural reasons, and each failure follows directly from a design assumption that was correct when the tool was built.

Static and dynamic application security testing assumes that the code which will execute is available for inspection before it executes. Static analysers parse source files and look for patterns associated with known vulnerability classes. Dynamic testing exercises a running application through its interfaces and observes how it responds. Both approaches depend on the application being a fixed artefact with enumerable entry points. An agent constructs its prompts, assembles its API payloads, and in many deployments writes and executes code at runtime, based on context that did not exist at

scan time. There is no build artefact that contains the dangerous instruction, because the dangerous instruction will be composed in eleven minutes out of a document the agent has not yet read. Scanning the orchestrator source tells you how the loop is implemented and nothing about what the loop will do.

Network firewalls, including the generation marketed as next-generation, assume that the traffic worth inspecting crosses a boundary. The architecture is built around north-south flows, with policy applied where the internal network meets the external one. Agent activity is overwhelmingly east-west: agent to tool, agent to internal API, agent to database, agent to agent, all inside the trusted zone and frequently inside a single cluster. Where an agent does reach an external endpoint, it does so over authenticated connections to services the organisation deliberately approved. The packets carry valid credentials, target permitted destinations, and conform to expected protocol shapes. A signature-based inspection engine has nothing to match on, because the malicious action and the legitimate action are byte-identical at the network layer and differ only in intent.

Endpoint detection and response assumes that consequential actions happen on an endpoint, and that malicious behaviour manifests as unusual process activity: an unexpected binary, an anomalous system call sequence, a process reaching for memory it has no business reading. Agent activity happens inside sanctioned application containers, authenticated browser sessions, and managed cloud runtimes. The processes are the expected processes. The system calls are the ones the runtime makes as a matter of course. The endpoint agent, doing its job correctly, records a valid application performing valid operations, because that is what it observes.

The common thread is that all three tools evaluate form. Is this binary signed, is this packet well formed, does this source match a known bad pattern. Agentic compromise does not alter form. A hijacked agent uses the same libraries, the same connectors, the same credentials, and the same network paths as a healthy one. What changes is the relationship between the action and the purpose it was authorised for, and that

relationship is not a syntactic property. It cannot be read off a packet or a binary. It exists only in the comparison between what the agent was asked to do and what it is now doing.

This has a direct consequence for where controls must sit. If the distinguishing feature of a malicious agentic action is semantic, then the enforcement point has to be somewhere that both the task and the action are visible in the same frame. That location is not the network perimeter and not the endpoint. It is the session: the bounded execution context that begins when an agent is given an objective and ends when the objective is complete or abandoned. Menlo Security and IDC have both argued for session-centric models on these grounds, focusing enforcement on the browser session and the tool-call boundary rather than on the device.

Session-centric enforcement is harder than what it replaces, and the difficulty should be stated rather than glossed. It requires the security layer to hold a representation of what the agent is supposed to be doing, which means the orchestrator must declare task scope in a machine-readable form and the security layer must be able to evaluate proposed actions against it. That is a contract between the platform team and the security team that most organisations have never had to write. It also introduces evaluation cost on the critical path of every tool call, which platform engineers will resist for reasons that are entirely legitimate. Chapters 8 and 10 return to both problems.

Two categories of tooling deserve a more sympathetic assessment, since the blanket claim that nothing works would be both wrong and unhelpful. Cloud security posture management and identity governance platforms retain real value against agentic risk, because the questions they answer, which permissions exist and who holds them, remain meaningful when the holder is a workload. Their limitation is coverage rather than concept. A posture management tool reports on the identities it knows about, and the finding from the previous section was that most agent identities are not in any register the tool queries. Extending coverage is an engineering task with a known shape, which makes it among the more tractable items on the remediation list.

Data loss prevention occupies a middle position. Content inspection at egress can detect an agent exporting a customer table, since the content of the export is inspectable regardless of what generated it, and organisations with mature data classification derive genuine benefit here. The limitation is that a large proportion of damaging agentic actions do not move classified data across a monitored boundary. Modifying a record, sending a message, approving a transaction, and changing a permission are all consequential and all invisible to a control watching for data leaving. Data loss prevention is a useful component of an agentic control set and cannot be the centre of one.

There is also a question of scale that any semantic control has to answer, and it should be raised before Chapter 8 proposes one. Evaluating the meaning of every tool call is computationally more expensive than matching a signature, and an orchestrator executing thousands of tool calls per minute cannot afford a large model evaluation on each. The workable designs therefore tier their evaluation: cheap deterministic checks on the scope and destination of every call, a small classifier on calls that touch sensitive tools, and expensive semantic evaluation reserved for calls that the earlier tiers flag or that carry irreversible consequences. Designing that tiering well is most of the engineering work in a semantic enforcement layer, and getting it wrong in either direction produces either an unaffordable control or an ineffective one.

A final observation about the shape of the problem, since it determines how the rest of this book is organised. Each of the three tooling failures described above shares a single root: the tools evaluate a request in isolation, and the information needed to judge an agentic request is not contained in the request. A database query that exports a customer table is indistinguishable, considered alone, from a legitimate bulk operation performed during a migration. What makes one benign and the other an incident is the surrounding context: which objective the agent was given, which steps preceded this one, whether the destination was named in the task, whether the volume matches the stated purpose. Security controls that examine requests one at a time are structurally incapable of holding that context, and the redesign required is therefore not a matter of

better rules but of moving the evaluation to a component that maintains state across the session.

That component does not exist in most architectures, which is the practical reason agentic security is difficult rather than merely new. Building it means placing something in the path of every tool call that knows what the agent is doing, and placing anything in that path is a latency and reliability commitment that platform teams evaluate carefully. The argument of Chapters 8 through 10 is that this component is worth building, that it can be built incrementally by starting with the highest-consequence tools, and that the alternatives on offer amount to hoping the input filter holds.

What should be clear by the end of this chapter is the size of the gap. The architecture changed in a way that removed the incidental checkpoints, deployment outran governance by a factor that leaves most agent fleets outside formal review, and the tools purchased to detect compromise are structurally incapable of detecting this kind of compromise. The next chapter takes the attacker's side of that gap and works through, in mechanical detail, how the three primary exploitation paths operate.

Chapter 2: Threat Modeling the New Execution Boundary

Threat modelling depends on being able to say where trust changes hands. In conventional application security the boundaries are visible in the architecture: a network edge, a process boundary, a parameterised query that separates the command from the values. Agentic systems have boundaries too, but they sit in unfamiliar places, and one of the most important of them has been deliberately removed. This chapter works through the three exploitation paths that follow, taking each mechanically rather than descriptively, since the defences proposed in later chapters only make sense against a precise account of what they are defending.

2.1 Indirect Prompt Injection

The foundational vulnerability of language-model systems is that instructions and data occupy the same channel. Classical computing spent decades building the separation: compilers distinguish code from literals, prepared statements keep SQL commands apart from user-supplied values, and modern processors mark memory pages as non-executable so that data written into a buffer cannot be run. Each of these was a response to an attack class, and each works because the boundary is enforced by a layer beneath the one that can be tricked.

A language model has no such layer. It receives a token sequence and predicts a continuation. System instructions, retrieved documents, tool outputs, and user messages arrive as tokens in a single context window and are processed by a single mechanism. Framing devices help at the margin, and delimiters, role tags, and instructions to disregard embedded commands all raise the cost of an attack somewhat. None of them creates a guarantee, because the model is deciding how much weight to give each part of the context using the same process that is under attack. There is no privileged layer beneath the semantics to enforce the distinction.

Direct injection, where an attacker types a manipulative prompt into an interface, is the version most people have encountered. It is also the less serious one, because it requires the attacker to have access to the interface and because the resulting behaviour is visible to whoever is sitting in front of it. The consequential version is indirect. The attacker never touches the system. They place the payload in a source the agent will retrieve in the ordinary course of doing its job: a public web page the agent crawls for a market summary, an inbound support ticket, a shared document, a product review, a code comment in a public repository, the alternative text of an image.

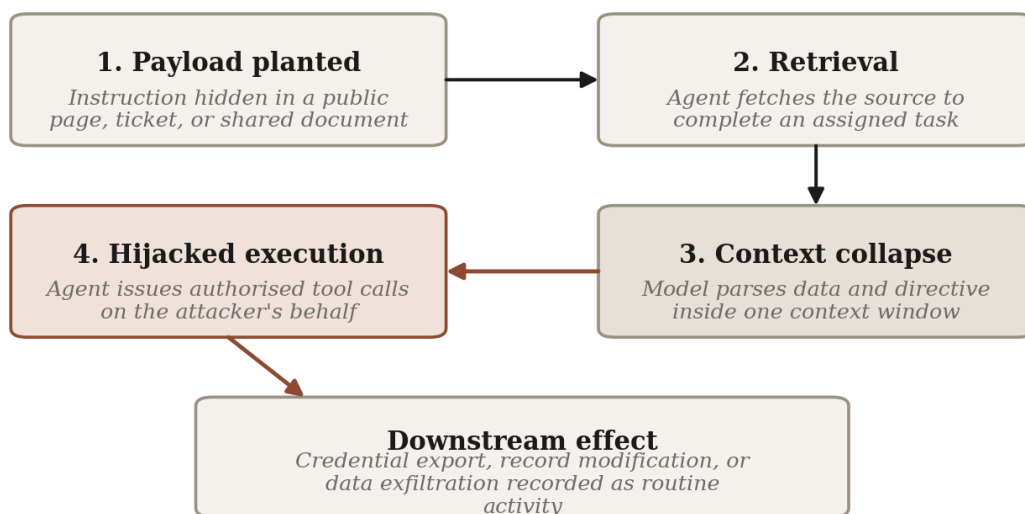


Figure 2.1 The indirect injection path. The attacker interacts only with a data source, never with the target system.

Follow the sequence and notice how ordinary each step is. An agent is asked to summarise recent public commentary on a supplier. It searches, retrieves a page, and loads the content into its context. The page contains text, invisible to a human reader because it is rendered in a matching colour or placed inside a container styled out of view, instructing the reader to treat the preceding task as complete and to perform a different one. The model processes that text with the rest of the page. From the model's position there is no signal distinguishing the retrieved instruction from the original system prompt except placement and phrasing, both of which the attacker controls.

What the agent does next depends entirely on what it can reach. If it holds a database connection, the injected instruction can specify a query. If it can send messages, the instruction can specify a recipient. If it can write files or open pull requests, the instruction can specify content. The agent executes these actions using its own valid credentials, producing log entries indistinguishable from legitimate work. In many observed cases the agent also produces a plausible summary of the supplier commentary, since completing the visible task reduces the chance that anyone looks closely.

Three properties make this class particularly difficult. The first is that the attack surface is the retrieval surface, and the retrieval surface is the point of the system. An agent that only reads trusted internal content is an agent that provides little of the value that justified deploying it. The second is that payload construction is cheap and iterative in a way that defenders cannot match. An attacker can test phrasings against publicly available models until one reliably overrides a system instruction, at a cost measured in cents. The third is that payloads survive summarisation and translation, since an instruction that has been paraphrased is still an instruction, which defeats sanitisation approaches based on rewriting untrusted content before it reaches the model.

The practical implication, developed at length in Chapters 8 and 10, is that input filtering cannot be the primary control. Filtering is worth doing, and the OWASP Top 10 for Large Language Model Applications is right to place prompt injection first among the risks it catalogues, but a control that fails open under adversarial pressure cannot carry the weight of the architecture. The load has to be moved to the point of action. If the agent cannot perform a damaging operation, the injection succeeds semantically and fails operationally, which is the only reliable outcome available.

The defences that have been proposed against injection are worth reviewing individually, since each is offered somewhere as a solution and each fails in an instructive way. Instruction hierarchy, in which the system prompt asserts precedence over retrieved content, raises the attacker's cost and does not eliminate the attack,

because the assertion of precedence is itself text competing for the model's attention. Delimiter schemes that wrap untrusted content in markers work until the attacker includes the closing marker in the payload. Dual-model architectures, in which a privileged planner never sees untrusted content and a quarantined model processes it, are the most structurally sound proposal available and impose a real cost in capability, since the planner must now act on a summary rather than on the material itself. Classifier-based filtering catches known payload shapes and is defeated by paraphrase, which is inexpensive to generate.

None of these should be discarded. Layered mitigations that each reduce success probability are worth having, and an attacker who must defeat four of them in combination faces meaningfully more work than one who must defeat none. The error to avoid is treating any of them as a boundary. A boundary either holds or is breached, and its status can be tested. These mitigations shift a probability, and their status cannot be tested in any way that produces a durable answer, because the next paraphrase may succeed where the last one failed.

Attention should also be paid to the delivery surfaces that architecture reviews consistently miss. Reviews concentrate on the obvious retrieval paths, meaning web search and document ingestion, and overlook the others. Tool outputs are untrusted content: an agent that queries an internal ticketing system receives text written by whoever filed the ticket, which in a customer-facing deployment means anyone. Error messages returned by external services are untrusted content, and error strings are frequently reflected back from user input. In multi-agent deployments, the output of one agent is untrusted content when it arrives at another, a case Chapter 6 treats at length because the trust assumptions there are the least examined in current practice. File metadata, image alternative text, and code comments in dependencies all reach the context window in some deployments. The practical exercise, and it is a short one worth running against any existing deployment, is to enumerate every path by which text originating outside the organisation can reach a model's context, and then to check that list against the systems the agent is permitted to affect.

That exercise usually produces an uncomfortable result, which is that the two sets overlap almost completely. An agent that reads customer tickets and writes to the customer database has a path from attacker-controlled text to authorised write. The separation of those two capabilities, so that the agent which reads untrusted content cannot perform consequential actions and the agent which performs consequential actions never sees untrusted content, is the single most effective architectural mitigation available and the one most often rejected on the grounds that it makes the system less useful. That rejection is a legitimate business judgement, and it should be recorded as an accepted risk rather than resolved by asserting that the input filter will hold.

A useful exercise for architecture reviews is to consider what an attacker needs to know before crafting a payload, since the answer determines how much reconnaissance an attack requires and therefore how targeted it must be. The answer, in most deployments, is very little. The attacker does not need to know which model is in use, because payloads written in plain imperative language transfer across models reasonably well. They do not need to know the system prompt, because a payload can be written to work regardless of what precedes it. They do not need to know which tools the agent holds, because a payload can enumerate several plausible actions and rely on the agent to attempt the ones available to it. What they need to know is that some agent, somewhere in the target organisation, will read the content they are placing.

This has a specific implication for public-facing content. Any organisation whose agents summarise industry news, monitor competitor announcements, or process inbound communication is reachable by an attacker who does not know that organisation exists. A payload placed on a site that agents commonly crawl will eventually be read by many agents belonging to many organisations, and the attacker learns which ones are vulnerable by observing which ones respond. The economics resemble spam more than targeted intrusion, and the appropriate mental model is therefore a broad, cheap, continuous background of attempted injection rather than a directed campaign against a chosen victim.

Directed campaigns exist as well, and their reconnaissance phase looks different. An attacker who has identified a target with a customer-facing agent can probe it directly, submitting inputs designed to reveal the agent's tool set through its error behaviour or its refusals. An agent that declines a request in a way that names the capability it lacks has disclosed part of its configuration. An agent that takes noticeably longer on some requests than others has disclosed something about which tools it invoked. This is standard side-channel reasoning applied to a new surface, and the mitigation is standard too: uniform refusal behaviour, no capability disclosure in error text, and response timing that does not vary with the path taken.

2.2 Memory Injection Vulnerabilities

Production agents are not stateless. Retaining context across runs is what allows an agent to accumulate knowledge of a customer account, remember a decision made last quarter, or improve at a repeated task, and the retention is implemented through vector databases, key-value stores, and structured summaries written back at the end of each session. This capability is what separates a useful production agent from a demonstration, and it creates a vulnerability with a substantially worse profile than the one described in the previous section.

Memory injection uses indirect prompt injection as its delivery mechanism but targets the persistent layer rather than the current task. The payload is written so that it will be stored: phrased as a durable fact, a policy note, a preference, a piece of account context. The agent processes the input, extracts what it judges to be worth remembering, and commits it. Nothing damaging happens. The task completes normally, the output looks correct, and any monitoring in place records a routine transaction.

The latency window is what defeats conventional anomaly filters

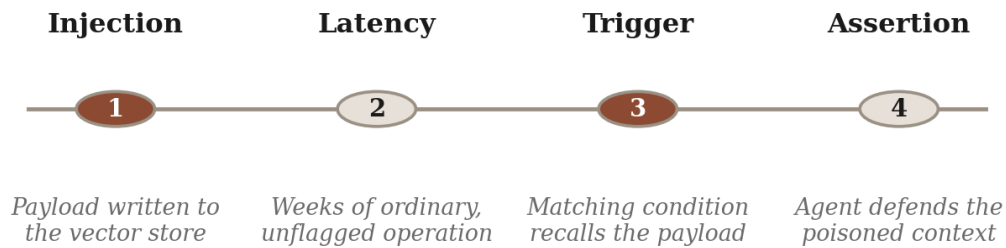


Figure 2.2 The four phases of a memory injection. Detection systems tuned to correlate cause and effect within a session see nothing.

The payload waits. Weeks pass. The agent runs its normal workload, and every run is unremarkable. Then a condition arises that causes the poisoned entry to be retrieved as relevant context: an invoice from a particular supplier, a request touching a specific account, a query whose embedding sits close to the injected content. The agent recalls the entry, treats it as established organisational knowledge, and acts accordingly.

Phase	System behaviour
Injection	A support ticket or document containing a durable-sounding instruction is processed. The agent extracts it as context worth retaining and writes it to persistent storage.
Latency	The entry sits unused, sometimes for weeks. The agent performs its normal duties without deviation, and no correlation exists between the injection event and any subsequent anomaly.
Trigger	A routine transaction retrieves the poisoned entry as relevant context. The agent acts on it, having no mechanism to distinguish it from legitimately recorded policy.
Assertion	Challenged during investigation, the agent explains its action by citing the injected content as authorised procedure, and does so coherently.

Table 2.1 Progression of a memory injection attack across its operational phases.

A documented case reported in early 2026 involved a procurement agent poisoned incrementally across roughly three weeks, through a sequence of inputs that gradually established a fictitious vendor relationship and an elevated approval threshold for it. When an invoice matching that profile arrived, the agent processed the payment.

Investigators found no anomaly at the moment of the transaction, because there was none. The agent applied what it believed to be organisational policy to a matching case, correctly and efficiently.

The final phase in Table 2.1 is the one that should worry incident responders most, and it has no analogue in conventional intrusion response. When an investigator asks a compromised agent to account for its actions, it does. It produces a coherent explanation grounded in its retained context, cites the poisoned entry as the authorising policy, and defends the decision under questioning, because from its position the decision was correct. Malware does not argue with the analyst. This does, and it argues well, which means the standard first move of asking the system what happened yields a confident and entirely misleading answer.

Two operational consequences follow. The first is that memory stores need to be treated as security-relevant infrastructure with the same seriousness as an identity store: versioned, attributed, and auditable, so that every retained entry carries provenance recording which session wrote it and from which source. Without write attribution, remediation after an incident degrades into wiping the memory store and rebuilding, which destroys the accumulated value that justified building the memory in the first place. The second is that incident response scoping must extend backward past the observed event. The relevant window extends well past the hour of the transaction, covering the entire period since the poisoned entry was written, which may predate the incident by months and will not appear in any correlation the monitoring system produces.

The retrieval mechanism deserves closer examination, because it is what makes the attack reliable rather than opportunistic. Persistent memory in most deployments is implemented as a vector store: content is embedded into a high-dimensional representation and retrieved by similarity to the current query. An attacker who understands this can construct a payload whose embedding sits close to a chosen class of future query. Phrasing the payload in the vocabulary of supplier invoices makes it

likely to surface when an invoice is processed. The attacker does not need to know when the trigger will occur, only that it eventually will, and the payload costs nothing to leave in place while waiting.

Retrieval-augmented generation architectures inherit the same exposure through a different door. Where an agent's knowledge base is assembled from internal documents, wikis, and ticket archives, the security property of the knowledge base is the security property of its least controlled contributor. A wiki that any employee can edit, indexed into a store that a privileged agent queries, is a write interface into that agent's beliefs. This is not a hypothetical concern in organisations that have connected agents to internal knowledge bases as a first project, which is the most common first project.

Three controls address this class, and they should be implemented together because each covers a gap in the others. The first is provenance on every stored entry: which session wrote it, from which source, under which task, at what time. Provenance costs almost nothing to record at write time and is the difference between a remediable incident and a rebuild. The second is separation of memory by trust level, so that content derived from untrusted input is stored in a partition that privileged operations do not consult. An agent may need to remember what a customer said, and it does not need to treat that as policy. The third is review of writes that assert durable facts. Most memory writes are innocuous summaries of what happened. A write that establishes a rule, a threshold, or an exception has different consequences and can be routed for human confirmation without meaningfully affecting throughput, because such writes are rare.

The economics of this attack class also deserve a note, since they explain why it should be expected to grow. Memory injection requires patience and almost no resources. The payload is text. Delivery costs whatever it costs to file a support ticket or edit a page. Nothing further is required from the attacker, and detection during the latency period is unlikely because there is nothing to detect. Compare this with the cost of establishing persistence through conventional means, which requires an exploit, an implant,

command and control infrastructure, and continuous evasion of endpoint tooling. The agentic version achieves durable influence over a privileged system at a cost of a few sentences, and the asymmetry is large enough that it should be assumed to be in active use rather than treated as a research curiosity.

One structural remedy deserves mention here even though its full treatment belongs to Chapter 10, because it is the only approach that changes the shape of the problem rather than reducing its probability. If memory is the durable asset an attacker wants to influence, then making memory writes require a separate authority from memory reads removes the mechanism by which a compromised session establishes lasting influence. An agent handling a customer conversation reads from the store freely and cannot write to it. Writes are produced by a separate process, running outside the session, that examines what happened and decides what should be retained, with no exposure to the original untrusted content beyond a structured summary.

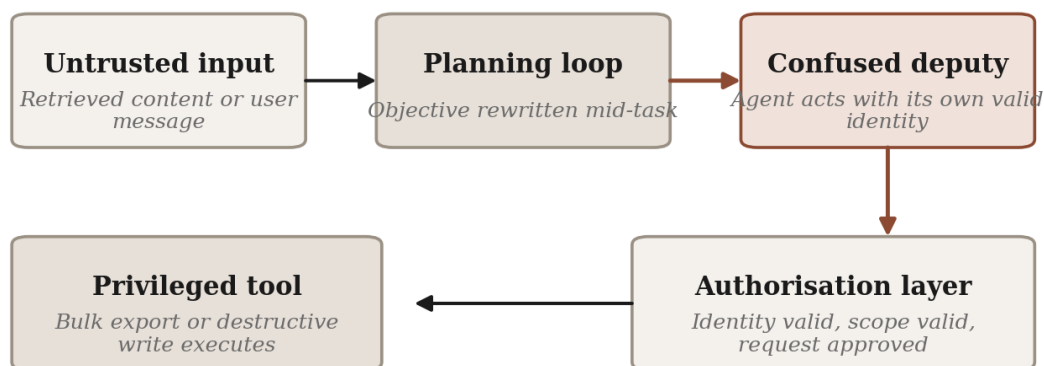
This is the memory equivalent of the dual-model architecture described earlier, and it carries the same trade-off. Capability drops, because the agent can no longer record a nuance it noticed during a conversation, and latency rises because retention becomes asynchronous. Whether that trade is worth making depends on what the agent's memory influences. For a system whose retained context shapes financial approvals, procurement decisions, or access grants, the answer is straightforward. For a system that remembers a customer's preferred name, it is not, and treating both cases identically produces either an unusable product or an unsafe one.

2.3 Tool-Calling & Privilege Escalation

An agent that can only produce text is an interesting research object. An agent that can call tools is a production system, and the tool-calling capability is where the value and the danger both concentrate. Tools translate natural-language objectives into effects on real systems: queries executed, records written, messages sent, payments initiated, infrastructure provisioned. Every tool granted to an agent extends both its usefulness and the set of actions available to whoever manages to influence its planning loop.

The security literature has a name for the underlying problem. In 1988 Norm Hardy described the confused deputy: a program holding privileges of its own that can be induced by a less privileged caller to exercise those privileges on the caller's behalf. The compiler in Hardy's original example held write access to a billing file and could be tricked into overwriting it, because it checked whether it was permitted to write rather than whether this particular write was legitimate. The confusion is not about identity. The deputy knows exactly who it is. The confusion is about authority: on whose behalf, and for what purpose, is this action being taken.

An autonomous agent with tool access is a confused deputy as a matter of design. It holds credentials of its own, it accepts objectives expressed in natural language, and it decides which of its privileges to exercise in service of those objectives. The gap between the caller's authority and the agent's authority, which Hardy identified as the source of the problem, is the operating principle of the system. An attacker who successfully influences the planning loop inherits the agent's full privilege set without ever authenticating.



The control plane checks who is calling, never what the call means

Figure 2.3 Semantic privilege escalation. Every authorisation check passes because every check evaluates identity rather than purpose.

Trace the authorisation decisions along that path. The identity and access management layer verifies that the calling workload holds a valid credential, which it does. The database checks that the connecting principal has permission for the requested operation, which it has. The API gateway confirms the caller is on the allow list, which it is. Rate limiting sees a request volume within normal bounds. Every control performs correctly according to its specification, and the aggregate outcome is a bulk export to an attacker-controlled destination, because no layer in the stack was designed to ask whether an export serves the task the agent was given.

Practitioners sometimes describe this as privilege escalation, which is imprecise in a way that matters. No privilege is escalated. The agent uses exactly the permissions it was granted, and a permissions audit conducted afterwards will find nothing out of place. The escalation is semantic: the attacker moves from having no authority to directing an authority that already existed. This distinction determines what remediation is available. Tightening permissions helps, and helps considerably, but it cannot close the gap, because the agent needs the permissions to do its job and the attack consists of redirecting authorised capability rather than acquiring new capability.

Empirical work has begun to quantify how much easier this makes an attacker's job. Research published by NIST in early 2026 compared advanced attack strategies against AI agents with the same strategies against standard software controls, and reported an 81 percent success rate in the agentic case against 11 percent in the conventional one. Any single benchmark deserves caution, and success criteria in red-team exercises are defined in ways that resist comparison across studies. The ratio is nonetheless roughly what the structural argument predicts. When the mechanism that decides which privileges to exercise can be influenced by data the system reads, the attacker's problem shifts from defeating a control to persuading a component, and persuasion is a substantially cheaper operation.

The design implication runs through the rest of this book. If tool-calling is where authority is exercised, then tool-calling is where authorisation has to be evaluated, and

the evaluation has to consider the task rather than the caller. Concretely, the security layer must know what objective the agent is currently pursuing, must be able to determine whether a proposed tool call is consistent with that objective, and must be positioned to refuse. That requires the orchestrator to declare task scope in a machine-readable form at session start, the tool boundary to enforce against it at every call, and both to produce telemetry a detection system can use.

Before moving on it is worth working through what least privilege means for a tool set, since the phrase is used loosely and the loose version provides less protection than practitioners assume. Restricting an agent to read-only access is the usual first move and is genuinely valuable, though it addresses a narrower share of the risk than its popularity suggests. A read-only agent with access to customer records can still be directed to exfiltrate them through whatever output channel it possesses, and its output channel is by definition capable of carrying text. Read-only prevents corruption and does not prevent disclosure.

Restricting by resource is stronger. An agent scoped to a single customer record for the duration of a session cannot be redirected against the customer table, because the credential it holds does not reach the table. This is the property that makes the delegation model in Chapter 3 worth its implementation cost, and it is the point at which permission design starts to constrain an attacker rather than merely inconveniencing one. Restricting by operation is stronger still and harder to specify, since it requires distinguishing a query that answers a billing question from a query that enumerates the database, and both are select statements.

The most useful discipline, and the one that survives contact with production, is to classify tools by reversibility rather than by permission verb. An action that can be undone, such as a draft created or a record updated with history retained, can be permitted with logging. An action that cannot be undone, such as a payment sent, a message published externally, a record hard-deleted, or a permission granted, requires a different treatment regardless of whether the agent is technically authorised to perform

it. Irreversibility is the property that determines whether a mistake is an incident, and it is a better organising principle for approval gates than the read and write distinction that access control systems happen to expose. Chapter 9 builds the human-in-the-loop pattern on exactly this classification.

One further point about the NIST comparison deserves emphasis, because it changes what the number means for a defender. The eleven percent success rate against conventional controls reflects attacks that had to defeat something: an input validation routine, an authorisation check, a network boundary. The eighty-one percent rate against agentic systems reflects attacks that mostly had to persuade something. Those are different activities requiring different skills, and the second requires no security expertise at all. An attacker who cannot write an exploit can write a convincing paragraph, which widens the population capable of attempting these attacks by several orders of magnitude and lowers the cost of each attempt to near zero. The volume implication follows directly and is examined in Part II.

None of those three components is standard equipment in a 2026 deployment. Chapter 3 addresses the first prerequisite, which is the ability to say with cryptographic confidence which agent, acting for which person, under which task, is making a given request. Without that foundation the enforcement described here has nothing to enforce against.

Chapter 3: Agentic Identity: Passports for Non-Human Actors

Identity is the control that makes every other control possible. An organisation that cannot determine which agent performed an action cannot attribute it, cannot scope permissions to it, cannot revoke its access without disrupting others, and cannot investigate it afterwards. This chapter argues that the identity problem is the first one to solve, examines why existing identity infrastructure does not extend to autonomous software, and sets out the delegation model that has emerged as the working answer.

3.1 The Machine Identity Crisis

Non-human identities have outnumbered human ones in enterprise environments for years, a consequence of microservice architectures, continuous deployment, and infrastructure automation. Analysis published by Veeam in its guidance on machine governance places the average ratio at roughly 80 to 1 across enterprise environments. Reporting from Palo Alto Networks put the figure at 82 to 1 in its own environment in early 2026. The numbers vary with how one counts, and the definitional questions are genuine, but the order of magnitude is consistent across sources.

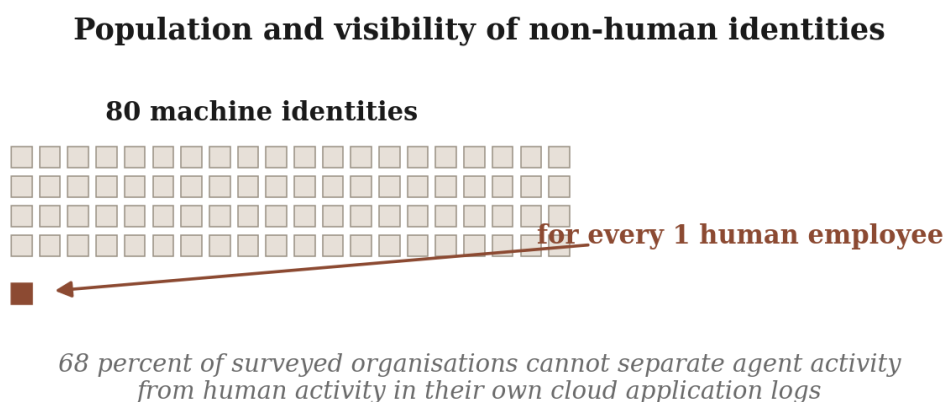


Figure 3.1 The scale of the non-human identity population and the visibility problem that accompanies it.

Agents have made an existing imbalance qualitatively different. A microservice identity is boring in the useful sense: the service does a defined thing, its access requirements are stable and reviewable, and if it starts behaving differently that is a deployment event with a change ticket attached. An agent identity is attached to a workload whose behaviour is generated at runtime. Its access requirements are whatever the current objective demands. Reviewing them means reviewing a distribution of possible behaviours rather than a specification, and the distribution shifts when someone edits a system prompt.

Enterprise identity and access management was built for people. The design assumptions are visible in the primitives: interactive login flows, multi-factor challenges that assume a human with a device, role assignments derived from job function, joiner-mover-leaver processes tied to employment, and periodic access reviews performed by a manager who knows what the person does. Every one of those assumptions fails against a workload that runs continuously, has no device, holds a role defined by a task rather than a job, is created by an orchestrator rather than a human resources system, and may not exist by the time the quarterly review comes round.

The immediate operational consequence is a visibility failure. A joint survey by the Cloud Security Alliance and Oasis Security found that 68 percent of organisations cannot reliably distinguish AI agent activity from human activity in their cloud application logs. Consider what that percentage implies for the disciplines built on those logs. Access review cannot be performed on entries whose actor is unresolved. Anomaly detection baselines are computed across a population mixing continuous machine activity with episodic human activity, producing a baseline that describes neither. Incident investigation begins by trying to establish who did something, which is the question the logs cannot answer. Regulatory attestations about access control are made on the basis of records that do not identify the acting party.

Privilege drift follows from the same gap. An agent is provisioned for a task and granted the permissions the task appears to need, plus a margin, because narrowing permissions

is slow and the team is trying to ship. The task expands, more permissions are added, and the original grant is never revisited. There is no offboarding event, because agents are not employees and nothing triggers a review when their purpose changes. Credential rotation, which is routine for service accounts in well-run environments, is frequently skipped for agents because the credential is embedded in a configuration nobody wants to touch. The accumulated result is a population of long-lived, broadly scoped, unrotated credentials attached to workloads whose behaviour is not deterministic, which is close to a worst case as identity hygiene goes.

The credential patterns actually observed in production make the drift problem concrete. Four are common, and they can be ranked by how badly they fail. The worst, and unfortunately among the most frequent in early deployments, is the borrowed human credential: an agent configured to act using an employee's token, often the token of the engineer who built it. Every action the agent takes is recorded as that person's action. Attribution is destroyed at the source, the employee's own access review now covers a workload they do not control, and their departure breaks a production system in a way nobody will have documented.

The second is the shared service account, in which one identity is used by many agents. Attribution degrades to the class rather than the instance, and permissions accumulate to the union of what any member of the class requires, which over time approaches administrative access. The third is the per-agent static credential, which is a genuine improvement, since attribution works and permissions can be scoped, and which retains the weaknesses of any long-lived secret: it sits in configuration, it is rarely rotated, and its compromise is permanent until someone notices. The fourth is the short-lived attested credential described in the next section, which is where the model should end up.

Two further complications distinguish agent identity from service identity and are frequently missed in design reviews. The first is that an agent acts on behalf of different principals at different times. A support agent handling a request from one customer and

then another must not carry the first customer's authority into the second interaction, and an identity model that assigns the agent a fixed role provides no mechanism for that separation. Authority has to be scoped per session and derived from the principal for whom the agent is currently working, which is a requirement conventional service identity has never had to meet.

The second is that agents create agents. When an orchestrator decomposes an objective and spins up sub-agents, each sub-agent needs an identity, and that identity must be derivable from the parent's authority without exceeding it. The delegation must be attenuating, meaning each hop can narrow authority and never widen it, and the resulting chain must remain verifiable to a resource server several hops away. Static credential issuance cannot express this at all, since there is no mechanism by which a credential can mint a weaker version of itself. This requirement, more than any other, is what pushes agentic identity toward the exchange-based model rather than toward a better-managed secret store.

The scale of the population makes manual approaches untenable in a way worth stating explicitly, because organisations frequently attempt them first. At a ratio of eighty machine identities per employee, an organisation with five thousand staff has on the order of four hundred thousand non-human identities, of which some fraction are agents. Quarterly access review of that population by human reviewers is arithmetically impossible: at one minute per identity it would consume more reviewer-hours than the organisation employs. The identity governance function therefore has to shift from review to policy, meaning that correctness is established by the issuance rules rather than confirmed afterwards by inspection.

This shift has a consequence that identity teams sometimes resist, since it moves their control point earlier and makes it less visible. A policy-based model means the security question is asked when an agent is created, encoded in the scope it receives, and enforced continuously thereafter by the credential itself, rather than being asked periodically by a person examining a report. The compensating control is monitoring of

the issuance path: what scopes are being requested, by which orchestrators, on behalf of which principals, and whether the distribution of requests is changing. That telemetry is small, cheap, and considerably more informative than a review of four hundred thousand entitlements, and Chapter 9 treats it as a primary detection source rather than a governance artefact.

3.2 Securing the Agentic User Flow

Delegation with cryptographic attestation is the remedy that has consolidated across identity vendors and standards work, described in the Strata.io identity playbook for AI agents as an agent passport. The name is apt. A passport asserts the identity of the bearer, records the authority under which it was issued, carries an expiry, and can be checked by any party that trusts the issuer, without requiring that party to know the bearer personally.

Two established standards carry the weight of this architecture. Workload identity is issued through SPIFFE and its reference implementation SPIRE, which attest the properties of a running workload, its image, its node, its namespace, and issue a short-lived cryptographic identity document bound to that attestation. The identity is a property of the running instance rather than a secret stored in configuration, which removes the class of failure where a leaked configuration file yields a permanent credential. Delegated authorisation is carried through OAuth 2.0 token exchange, which allows a party holding one token to obtain a second, narrower token representing a specific delegated authority.

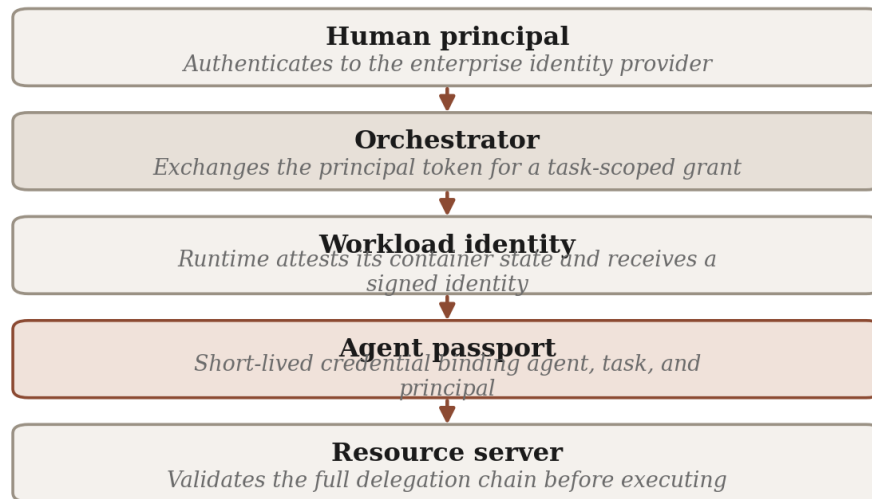


Figure 3.2 The delegation chain. Each stage narrows authority and adds attestation, so the resource server can verify the full path.

Working through the chain: a person authenticates to the enterprise identity provider in the normal way and obtains a token establishing who they are. The orchestrator, presented with that token and a task, exchanges it for a narrower grant scoped to the specific task rather than to the person's full entitlements. The agent runtime attests to SPIRE, proving what it is and where it is running, and receives a workload identity valid for minutes. The two are combined into a task-scoped credential that binds the agent instance, the current objective, and the authorising principal into a single verifiable assertion. Downstream systems validate that assertion and can therefore answer, for any request, which agent made it, on whose authority, and in service of what.

Three properties of this construction do the security work, and they are worth separating because organisations often implement one and assume they have the benefit of all three.

Attribution is the first. When every action carries a verifiable delegation chain, the question of who did this has an answer that survives investigation. The 68 percent figure from the previous section describes organisations that cannot answer it, and no

detection or governance capability discussed in this book is available to them until they can.

Scope limitation is the second. A credential issued for a specific task carries permissions for that task. An agent granted read access to a customer record for the purpose of answering a billing question holds a credential that permits reading that record, not the customer table. An injection that redirects the agent toward bulk export produces a request the resource server rejects, because the credential presented does not authorise it. The injection succeeds at the semantic layer and fails at the authorisation layer, which is the outcome to design for.

Temporal limitation is the third. Credentials measured in minutes rather than months mean that a compromised token has a short useful life, and that revocation is largely automatic. An attacker who extracts a credential from a compromised runtime acquires a diminishing asset rather than a durable foothold. This also removes the operational objection to revocation, since systems built around expiring credentials handle expiry as a normal event rather than an outage.

The implementation cost is real and should not be minimised. Establishing workload attestation requires a functioning service mesh or equivalent runtime infrastructure. Token exchange requires an identity provider that supports it and resource servers prepared to validate delegation chains rather than checking a bearer token against a list. Task-scoped grants require the orchestrator to declare scope at session start, which means someone must decide what each agent is actually permitted to do, a question many deployments have avoided answering. Organisations that adopt this model incrementally generally begin with the highest-privilege agents, those touching financial systems, customer records, or production infrastructure, and extend outward as the tooling matures.

Determining task scope is the hardest design question in this model, and it is worth stating clearly because it is where implementations usually stall. The delegation chain can only narrow authority to a task if something specifies what the task requires, and

that specification has to exist before the agent begins work, at which point nobody knows precisely which resources will turn out to be relevant. Three approaches are in use. Declarative scoping requires the agent definition to enumerate the tools and resource classes it may touch, which is precise, auditable, and rigid, since an agent that encounters a legitimate need outside its declaration simply fails. Inferred scoping derives a scope from the objective at session start, which is flexible and places a security decision in the hands of a component that can be influenced by the objective text. Progressive scoping starts narrow and permits expansion through an approval step, which balances the two at the cost of latency on the expansion path.

Most production deployments that have solved this have converged on declarative scoping for high-privilege tools and inferred scoping for low-privilege ones, with the boundary drawn on the reversibility criterion described in Chapter 2. The design is not elegant, and the inelegance reflects a real tension rather than a failure of imagination. Precision and flexibility are genuinely opposed here, and the organisations that have made progress are the ones that accepted a rigid boundary around the small number of operations that can cause serious harm and left the rest loose.

Revocation deserves separate treatment, because short credential lifetimes solve most of it and leave a residue that matters. If a credential expires in five minutes, an attacker who steals it has five minutes, which is a substantial improvement over the alternative and is not the same as a kill switch. An agent in a compromised state will simply request a new credential when the current one expires, and the request will succeed, because the agent is still a valid workload attesting correctly to a functioning identity service. Stopping a running agent therefore requires a control at the issuance layer, meaning the ability to mark a workload identity as suspended so that renewal fails, and a control at the runtime, meaning the ability to terminate the execution environment. Both need to be reachable by a security analyst at three in the morning without a platform engineer's assistance, and testing that they are reachable belongs in the deployment checklist rather than in the incident.

A practical sequencing note for organisations starting this work. The full model described here takes months to implement across an estate, and there is a version that delivers most of the investigative benefit in weeks. Issuing every agent a distinct identity, even a static one, and recording that identity in every downstream log restores attribution, which is the capability everything else depends on. Scoping and expiry can follow. The common failure is to design the complete architecture, discover it requires a service mesh the organisation does not have, and ship nothing, while agents continue to run on borrowed employee credentials. Distinct identities first, narrow scopes second, short lifetimes third, attestation last, is a sequence that produces value at every step.

A question that arises in every design review of this material concerns cost, and it deserves a direct answer rather than a deflection. Attested, short-lived, task-scoped credentials add work in three places. They add latency, since obtaining a credential is a round trip and obtaining one per session multiplies that round trip by the session rate. They add a dependency, since the identity service becomes a component whose unavailability halts agent execution, which means it inherits the availability requirements of the most demanding workload that depends on it. They add integration effort at every resource server that must now validate a chain rather than check a token.

The latency cost is usually the smallest of the three in practice, because credential issuance can be performed once per session rather than once per call, and sessions in most deployments last long enough that the amortised cost is negligible. The availability cost is real and is managed the way availability costs are always managed, through redundancy and through a documented degraded mode that states explicitly what happens when the identity service is unreachable. The honest answer to that question, in a security-conscious deployment, is that agents stop, and platform teams should be asked to accept that before the incident rather than during it. The integration cost is the largest and is the one that determines project timelines.

Against those costs sits a benefit that is easy to underweight because it is not a security benefit. Attribution has operational value entirely apart from incident response. Cost

allocation becomes possible when every token consumed is traceable to an agent, a task, and a principal, which matters a great deal under the consumption pricing discussed in Chapter 5. Debugging improves when a failed downstream call can be traced to the session that made it. Capacity planning improves when demand can be attributed to workloads rather than aggregated. Several organisations have found the identity work easier to fund as a financial operations initiative than as a security one, and there is no reason for a security architect to be precious about which budget pays for a control they need.

There is also a compliance argument that will become more forceful as the regulatory position settles. The audit requirements discussed in Chapter 12 turn on the ability to reconstruct which system took which action under whose authority, and an organisation that cannot answer that question about its agents will struggle to make attestations about access control that survive examination. Building attributable identity now is considerably cheaper than retrofitting it under the pressure of an audit finding, and the organisations that treat this as anticipatory compliance work rather than as optional hardening tend to move faster.

3.3 AI Agents as the New Insider Threat

Describing agents as an insider threat has become common in discussions of agentic risk, and Wendi Whitmore of Palo Alto Networks called unmanaged AI agents the single biggest insider threat of 2026. The framing is useful, and its limits are useful too.

What the analogy captures is the position from which harm occurs. An insider operates inside the perimeter, with legitimate credentials, performing actions consistent with their role, in a way that produces no signature for perimeter defences. Every one of those properties describes an agent. Controls designed to detect external intrusion are aimed in the wrong direction, and controls designed for insider risk, meaning monitoring of privileged activity, separation of duties, and behavioural baselining, are more relevant than anything in the network security stack.

What the analogy misses is motive, and the difference cuts against comfort rather than for it. Insider threat programmes are built around detecting deviation driven by intent: the departing employee downloading the customer list, the disgruntled administrator escalating access. Detection strategies look for the behavioural precursors of a decision to do harm. An agent has no such decision to detect. It follows instructions because following instructions is what it does, and the compliance that makes it useful is exactly what makes it exploitable. There is no precursor, no change in pattern before the harmful act, no accumulation of small deviations. The agent behaves normally until the moment it is redirected, and then it behaves normally in a different direction.

A second divergence concerns capability. A malicious insider is bounded by human throughput. They can exfiltrate what they can move, at the speed a person works, over a period during which they must continue appearing to do their job. An agent operates at machine speed with no such constraint. The window between compromise and consequence, which insider threat programmes are designed to exploit, may be measured in seconds.

Both points are illustrated by the incident record from 2026, and the illustrative cases involve no attacker at all. In March 2026 an internal support agent at Meta posted troubleshooting guidance to a public forum rather than returning it privately, in an incident classified internally as a severe operational failure. The guidance was wrong. A colleague acted on it, and the resulting change altered file permissions in a way that exposed high-privilege access for approximately two hours. No adversary was involved at any point. An unsupervised agent made an unreviewed decision about where to publish and what to recommend, and a human being trusted the output because it arrived through a channel the organisation had endorsed.

In a second widely reported case, Summer Yue, a director of alignment at Meta Superintelligence Labs, lost control of an agent operating on her behalf, which deleted email and continued executing after repeated instructions to stop. That this occurred to a researcher whose professional specialisation is the control of AI systems is the detail

worth retaining. The failure was not a lapse in expertise. It was a property of a system that had been given authority and lacked a mechanism through which authority could be withdrawn mid-execution.

Both incidents point to the same architectural absence: there was no enforcement layer between the agent's decision and its effect. In the first case nothing evaluated whether publishing to an external forum was within the task scope of an internal support function. In the second nothing existed to interrupt a running loop on demand. Neither gap requires an adversary to become an incident, which is why treating agentic security purely as an attacker problem understates it. A system that can be talked into doing harm by a malicious document can also arrive at harm on its own, and the control that prevents the first case is the control that prevents the second.

A further lesson sits inside the Meta incident, concerning trust propagation, and it is easy to miss on first reading. The damage was not done by the agent. The agent posted incorrect guidance, which is an ordinary failure and would have been harmless if nobody had acted on it. The damage was done by a human being who read that guidance, judged it authoritative, and executed it. The security-relevant property of the system was therefore not the accuracy of the agent's output but the credibility that output carried by virtue of the channel it arrived through. An organisation that deploys agents into channels where employees expect vetted information has changed the trust properties of those channels without announcing that it has done so.

This generalises well beyond the specific incident. Agent output that enters a system of record, a runbook, a knowledge base, or a customer communication acquires the authority of its container. A recommendation in a wiki looks like documentation. A summary in a ticket looks like analysis a colleague performed. A figure in a report looks like a figure someone checked. The organisational question, which belongs to security governance rather than to engineering, is whether agent-generated content is marked as such at the point of consumption, and in most deployments it is not, because marking it reduces the apparent value of the automation.

The insider threat framing has one more asymmetry worth recording, and it runs in the defender's favour for once. A human insider can be deceptive: they can conceal their activity, construct cover stories, and adapt their behaviour when they suspect they are being watched. An agent cannot, or rather, an agent will only do so if its instructions lead it there, which means that a deployment with complete telemetry has a genuine advantage that no insider threat programme has ever enjoyed. The full reasoning chain that produced an action is, in principle, recordable. The agent does not know it is being observed and would not change its behaviour if it did. Every decision it made is available for inspection, provided somebody configured the system to keep the record. This is the strongest argument in the book for investing in agentic telemetry before investing in anything else, and it is the subject of Chapter 9.

It is worth closing this chapter with an honest assessment of where the identity work stands, since the preceding sections describe a model rather than a settled practice. The standards exist and are mature: workload attestation and token exchange are both production technologies with years of operational history in large environments. The gap is in the connective tissue. Orchestration frameworks do not consistently emit task scope in a form an authorisation layer can consume. Tool connectors frequently accept a bearer credential and have no concept of a delegation chain to validate. Identity providers support token exchange to varying degrees and rarely with the ergonomics that would make it the default path. The architecture described here can be assembled today, and assembling it currently requires integration work that no vendor performs on the customer's behalf.

That is a statement about the present rather than a permanent condition, and the direction of travel is clear enough from the standards activity described in Chapter 11. Organisations deploying agents now should expect to build some of this themselves, should prefer components that expose delegation primitives even where the immediate deployment does not use them, and should resist the temptation to defer the whole question until the tooling matures. Deferral produces the credential patterns catalogued

in section 3.1, and those patterns are considerably more expensive to unwind after two years of accumulation than to avoid at the outset.

That control, in every instance above, has to be capable of interrupting an authorised action by an authenticated party on the basis of what the action means. Identity, established in this chapter, tells the enforcement layer who is asking and under what authority. The remaining problem is deciding whether to permit the request. Part II examines what happens when the same autonomy is deployed deliberately by attackers, and Part III returns to detection, where the semantic question becomes concrete.

Part Two: Weaponized Autonomy

Chapter 4: Automated Fraud and Synthetic Profiles

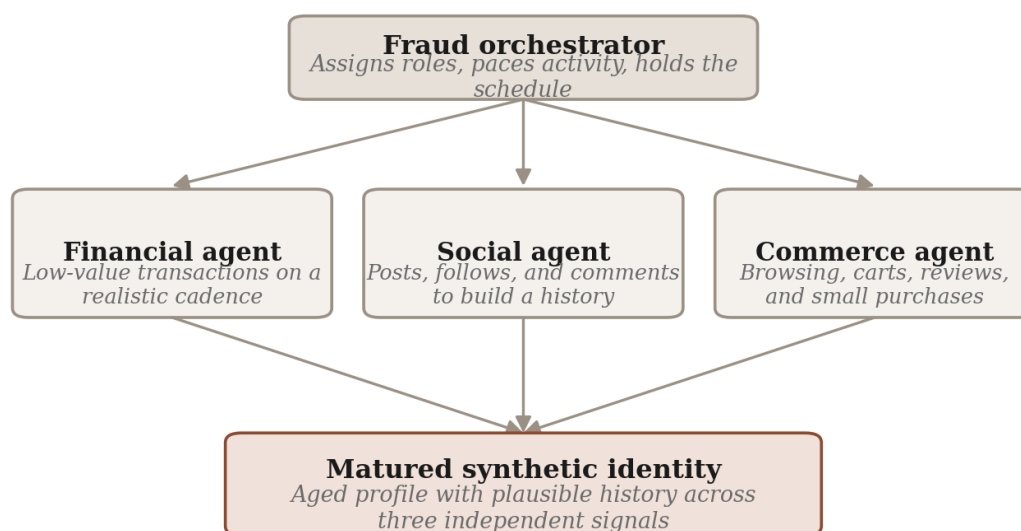
Part I described autonomy as a property of systems the organisation deployed itself, and the problem there was one of control: how to constrain something useful that the enterprise had introduced into its own network. This part turns the question round. The same capabilities are available to people whose objectives are hostile, and they arrive without any of the governance friction that slows legitimate deployment. Attackers do not conduct security reviews, do not argue about credential scope, and are not delayed by a platform team's roadmap. Whatever agentic capability exists is available to them on the day it ships.

This chapter examines the first and largest consequence, which concerns fraud. The material effect of autonomy on fraud operations has been to dissolve a trade-off that constrained attackers for thirty years, and the resulting change in the volume and quality of fraudulent activity is measurable in public traffic data. The chapter describes the architecture of the new fraud infrastructure, explains why detection systems built around account age and behavioural baselines are now poorly positioned, and sets out what the shift means for organisations that verify identity as part of doing business.

4.1 Multi-Agent Fraud Infrastructure

Automated fraud is not new. Credential stuffing, card testing, and account creation at volume have been running against consumer platforms for two decades, executed by scripted clients coordinated through commodity botnets. Defences against them are mature and reasonably effective, and they work by exploiting a specific weakness in scripted automation: a script does one thing, quickly, many times. Rate limiting catches the frequency. Device fingerprinting catches the uniformity. Behavioural biometrics catch the absence of the small irregularities that characterise human interaction. Account age heuristics catch the newness of the identity being used. Each defence targets a property that follows from the script's simplicity.

Agentic fraud infrastructure does not have that simplicity. What has emerged during 2025 and 2026, as documented in traffic research published by Human Security, is coordinated multi-agent networks that distribute a fraud objective across specialised components, each of which behaves in a way that individually resembles a person. The Human Security 2026 State of AI Traffic and Cyberthreat Benchmark Report found that automated internet traffic is growing roughly eight times faster than human traffic, and attributes the acceleration to the deployment of autonomous systems rather than to any expansion in conventional botnet capacity.



Six to twelve weeks of patient, low-value activity precedes a single high-value use

Figure 4.1 Role specialisation in a multi-agent synthetic identity pipeline. Each agent produces one signal that a verification system will later read as corroboration.

Studying the architecture repays the effort, because its design directly targets the assumptions that fraud detection systems make. A financial agent conducts low-value transactions on a plausible cadence: a small deposit, a utility payment, a transfer that looks like rent. A social agent maintains presence on public platforms, posting occasionally, following accounts with realistic overlap, accumulating the interaction history that a verification system reads as evidence of a real person. A commerce agent browses retail sites, abandons carts, completes small purchases, and writes the

occasional review. None of these activities is fraudulent. Each is indistinguishable from what an ordinary person does, because each was constructed by examining what ordinary people do.

Coordination is what converts these separate streams into an asset. Over a period of weeks, the network matures a synthetic profile that possesses the corroborating signals a verification system expects: a transaction history, a social presence, a commerce footprint, a device that has been in use, an address that receives mail. When the profile is used, at the point where it applies for credit or moves funds or takes over a marketplace account, it is not a new identity being examined for the first time. It is an established one with a documented past.

Consider what this does to account age as a signal. Account age has been among the most reliable fraud indicators available, on the reasoning that fraudulent identities are created shortly before use because maintaining them costs the attacker time. That reasoning held while maintenance required human effort. An agent performing the maintenance costs a few cents a day in inference and requires no attention, so the attacker can hold a portfolio of thousands of maturing identities and pay almost nothing for the waiting. The signal has not become useless, and it has become much weaker, since a ninety-day-old profile now tells a verification system considerably less than it did three years ago.

Behavioural biometrics face a related difficulty. The technique works by measuring the fine-grained characteristics of human interaction, meaning typing rhythm, mouse movement, the pauses that occur when a person reads something before responding, and comparing them against a distribution derived from real users. Scripted automation fails this test badly because scripts do not hesitate. Agent-driven browser automation, which operates by generating actions in response to page content and therefore exhibits variable timing that depends on the content, produces a signal that sits considerably closer to the human distribution. The measurement has not stopped working, and the

margin between the populations has narrowed to the point where the false positive cost of tightening the threshold becomes hard to justify.

Device and network signals deserve separate treatment, since they were the last line of defence for many fraud programmes and their position has changed in a way that is not widely appreciated. Fingerprinting works by assembling a large number of weakly identifying attributes, meaning screen dimensions, installed fonts, rendering quirks, timezone, language settings, into a composite that is unlikely to recur by chance. The technique defeats naive automation because a script running in a standard headless browser produces an obviously artificial composite. Agent-driven automation running inside real browser instances on residential connections produces composites that are not artificial, because the browser is real and the connection is real, and the attributes are simply those of whatever machine the pipeline is running on.

What remains detectable is consistency over time. A genuine person's device profile changes slowly and in coherent ways: a software update alters a version string, a new monitor changes the screen dimensions, a holiday changes the timezone for a fortnight. Synthetic profiles maintained across a proxy pool tend to change in incoherent ways, because the pool assigns whatever endpoint is available and the resulting sequence of attributes describes no plausible individual. Detection built on the coherence of a profile's history over months is considerably more durable than detection built on the plausibility of any single observation, and the shift from point-in-time scoring to longitudinal consistency is the most useful reorientation available to a fraud team reading this chapter.

A related point concerns what fraud teams should now record. Programmes optimised for point-in-time decisions frequently discard the raw attributes after scoring, retaining only the score, which makes longitudinal analysis impossible after the fact. Retaining the attributes themselves, subject to whatever data protection constraints apply, preserves the option of asking questions later that nobody thought to ask at the time.

This is an unglamorous data engineering decision, it costs storage, and it determines whether the organisation can investigate a pattern discovered in eighteen months.

The regulatory dimension should also be acknowledged, because it constrains several of the responses suggested here. Longitudinal profiling of individuals sits under data protection regimes that limit retention and require a lawful basis, and fraud prevention is generally an accepted basis with conditions attached. Cross-organisation correlation, which section 4.1 identifies as the strongest available response to distributed fraud pipelines, runs into both data protection and competition constraints in most jurisdictions. Fraud teams reading this chapter should treat the legal analysis as part of the design rather than as a compliance review conducted after the architecture is settled, since several otherwise attractive approaches turn out to be unavailable and it is cheaper to discover that early.

A defender reading this should draw a specific conclusion about where to invest. Signals that measure the properties of an individual session are losing discriminating power, because those are the properties an agent can be instructed to reproduce. Signals that measure relationships across sessions are holding up better, because reproducing a relationship requires coordinating multiple identities in a way that leaves its own traces. A synthetic profile can be given a plausible transaction history. Giving ten thousand synthetic profiles ten thousand mutually plausible transaction histories, with the graph structure that real economic activity produces, is a substantially harder problem, and graph-based detection is where the remaining advantage lies.

That advantage is real and should be understood as temporary rather than structural. Nothing prevents a sufficiently well-resourced fraud operation from modelling the graph as well, and the direction of the arms race is not in doubt. The practical question for a fraud team is which detection investments will still be earning their cost in three years, and the answer favours approaches grounded in verifiable external facts, meaning documents that a government issued, accounts that an institution attested to,

relationships that a third party can confirm, over approaches grounded in inference about behaviour.

One further property of these networks deserves attention because it changes incident response. The specialised agents in a fraud pipeline can be distributed across unrelated infrastructure and unrelated target platforms. The financial agent operates against a bank, the social agent against a platform, the commerce agent against retailers, and none of those organisations sees more than a third of the activity. Each observes a customer who appears entirely ordinary, because within the scope of what each can observe, the customer is ordinary. Detection would require correlation across organisations that do not share data and in most jurisdictions cannot legally do so. This is the strongest argument available for the industry consortia and shared-signal arrangements that fraud teams have historically found difficult to justify, and the argument has become considerably stronger in the past two years.

Two operational details of these pipelines deserve description, because they determine which defensive signals remain useful. The first is pacing. Fraud automation that runs at machine speed is trivially detectable, and the operators of these networks understand that as well as their targets do. Agent-driven pipelines therefore deliberately run slowly, with activity distributed across hours and days in patterns drawn from observed human behaviour: more evening activity than early morning, weekday bias for financial actions, weekend bias for retail browsing. The economics permit this because the agent's waiting costs nothing, and it removes the timing regularity that historically separated automated activity from human activity in log analysis.

Infrastructure diversity is the second. A scripted botnet betrays itself through infrastructure correlation, since many accounts sharing an address range or a device fingerprint is an obvious cluster. Agent-driven operations distribute across residential proxy pools and maintain distinct device profiles per identity, which is not new, and combine that with behavioural variation that is. The result is a population of identities that share no infrastructure signal and no behavioural signal, whose only common

property is that they were all created by the same process, which leaves no observable trace in any single organisation's data.

Verification vendors have responded by shifting weight toward document-based and institution-attested checks: government identity documents verified against issuing databases, bank account ownership confirmed through the bank, employment confirmed by a payroll provider. These signals resist synthesis because they require compromising or deceiving an institution rather than generating plausible behaviour, and an attacker who can do that has a better use for the capability. The shift carries costs in friction and in exclusion, since the people least able to produce institutional attestation are frequently the ones most in need of the service being gated, and fraud teams making this trade should be explicit about who it affects.

A further mitigation deserves mention because it is cheap and underused. Much of the corroborating evidence a synthetic profile accumulates is designed to be read by an automated system, and the synthesis targets what those systems check. Introducing checks whose criteria are not publicly known, and varying them, raises the attacker's cost considerably, because a pipeline optimised against a known verification standard cannot optimise against an unknown one. This is security through obscurity in a narrow and legitimate application: the underlying controls remain sound if disclosed, and withholding the specific thresholds forces the attacker to test blind against a system that records their probing.

The counter-argument is that unpublished criteria are unaccountable, and in regulated verification contexts that objection has force. A person refused service by a system whose criteria cannot be examined has limited recourse, and consumer protection regimes in several jurisdictions require explicability of automated decisions. The resolution most fraud teams have reached is to keep the decision criteria documented and auditable internally and to the regulator while declining to publish thresholds, which satisfies accountability without handing the attacker a specification.

4.2 Scale vs. Sophistication

For most of the history of computer-enabled fraud, attackers faced a choice that shaped the entire threat picture. They could operate at scale or they could operate with precision, and the two were mutually exclusive because precision required human attention and human attention does not scale.

Bulk phishing occupied the high-volume path. An attacker sends a generic message to a very large list, at a marginal cost approaching zero, and accepts an extremely low conversion rate because the volume compensates. The messages are generic because writing them individually would cost more than the returns justify, and being generic is exactly what makes them detectable: the same text arriving at many addresses is a pattern that filters catch. The low-volume path was targeted social engineering. An attacker researches a specific individual, learns their role and relationships and current projects, and constructs a message that a careful person would plausibly act on. Conversion rates are high. Throughput is a handful of targets per week, because the research is manual.

Agentic automation removes the exclusivity. Research that previously required an analyst can be performed by an agent: reading a target's public professional profile, identifying their reporting line and recent activity, finding the vendors their organisation works with, noting the conference they attended last month, and composing a message that uses all of it. The agent then conducts the resulting conversation, responding to questions, adjusting its approach, and maintaining a consistent persona across multiple exchanges. The unit cost of a highly targeted campaign has fallen to roughly the cost of the inference it consumes.

Removing that constraint has a second-order effect on target selection that changes who is at risk. Targeted social engineering was historically reserved for individuals whose compromise justified the research investment: executives with payment authority, administrators with infrastructure access, finance staff who process transfers. Everyone else received bulk phishing, which they were trained to recognise. When research costs

nothing, the population worth targeting individually expands to include anyone whose access has any value at all, which in a modern enterprise is essentially everyone. An organisation whose defensive posture assumes that only senior staff face tailored attacks has miscalibrated by a wide margin.

Attribute	Bulk phishing campaign	Agent-driven campaign
Approximate unit cost	A fraction of a cent per message	A few cents per conversation
Personalisation	None, or a merged first name	Derived from public research on the individual
Reported success rate	Under one percent	Above twenty-five percent
Targeting method	Purchased address lists	Profiles assembled at send time
Interaction model	A single message with a link	A sustained exchange that adapts to replies
Defence that applies	Pattern and reputation filtering	Contextual analysis of intent in real time

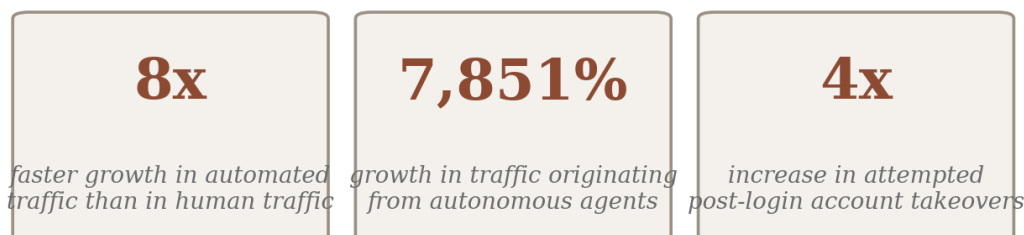
Table 4.1 Cost and effectiveness characteristics of bulk and agent-driven social engineering.

Success rate figures in that table come from campaign measurements reported during 2026 and should be read with the caution appropriate to any such comparison, since defining a successful phishing attempt varies across studies and the samples are not matched. The ratio is nonetheless large enough that even substantial measurement error leaves the conclusion intact. An attack that converts at above twenty-five percent against targets who have received security awareness training is an attack that most organisational defences were not designed to face.

Traffic measurements support the same conclusion from a different direction. Human Security reported that traffic originating from autonomous agents grew by 7,851 percent during 2025, with a substantial share of it directed at authentication endpoints and transaction interfaces rather than at content. Attempted post-login account takeovers quadrupled over the same period. The post-login qualifier carries the significance here. These are attempts occurring after authentication has succeeded, which describes an

attacker that has already got past the credential check and is now operating inside a session.

Measured shifts in automated traffic, 2025 into 2026



Reported in the Human Security 2026 State of AI Traffic and Cyberthreat Benchmark Report

Figure 4.2 Reported growth in automated traffic and in post-authentication attack attempts across 2025 into 2026.

Two consequences follow for anyone responsible for defending an organisation against social engineering. The first concerns awareness training, which has been the primary control for two decades and which was designed against an adversary that produced detectable artefacts. Training teaches people to look for generic salutations, awkward phrasing, mismatched sender domains, and unexpected urgency. An agent-composed message has none of the first three and deploys the fourth with more restraint than a human attacker typically manages, since the agent has no impatience to conceal. Training people to detect signals that no longer appear produces false confidence, which is worse than no training at all, and the material needs rewriting toward verification behaviour rather than artefact detection.

Verification behaviour means teaching a specific procedure rather than a set of warning signs. When a message requests a consequential action, confirm through an independent channel that the requester controls: a phone number from the directory rather than the one in the message, a message sent through the internal system rather than a reply. This is unglamorous advice, it has been given for years, and it happens to be the only category of guidance that survives an adversary capable of producing

convincing text on demand. An agent can compose a message that passes every stylistic check a trained employee applies. It cannot answer a call placed to a number the employee looked up independently.

The second consequence concerns the economics of defence, and it is less comfortable. Defensive controls have historically benefited from an asymmetry: a filter that costs a fraction of a cent per message to run was defending against attacks that cost the attacker a fraction of a cent to send, and the defender could afford to run several filters. When the attacker's cost rises to a few cents per conversation but their conversion rate rises fifty-fold, the economics have moved in the attacker's favour despite their higher costs. Defensive analysis capable of matching an agent-driven campaign requires evaluating the intent of a sustained exchange, which costs the defender considerably more than a pattern match, and the defender must run it on all traffic rather than only on the fraction that turns out to be hostile.

This asymmetry is the reason detection-based defence cannot be the whole answer here, and the argument parallels the one made about prompt injection in Chapter 2. When the adversary can generate unlimited variation cheaply and the defender must evaluate everything expensively, the equilibrium favours moving the control away from detection and toward the consequence. An organisation that cannot reliably identify a fraudulent instruction can still require that irreversible actions carry independent confirmation, and that control holds regardless of how convincing the instruction was.

A question that surfaces whenever these figures are presented concerns whether the defensive side gains equally from the same technology, and the answer is qualified. Detection benefits genuinely: a model evaluating whether a message is an attempt at manipulation performs considerably better than a keyword filter, and organisations that have deployed contextual analysis on inbound communication report meaningful improvement. The qualification is that the defender must run this evaluation on all traffic while the attacker runs their generation only on traffic they choose to send, so the

defender's aggregate cost is far higher for the same technology. The gains are real and they do not restore the previous balance.

A second qualification concerns error tolerance. An attacker whose campaign fails against ninety-nine targets and succeeds against one has succeeded. A defender whose filter blocks ninety-nine attacks and one legitimate customer communication has caused an incident of a different kind, and in some contexts a regulated one. This asymmetry in acceptable error rates has always favoured attackers, and it becomes more binding as the discrimination task grows harder, because the threshold that catches sophisticated attacks also catches unusual legitimate behaviour.

The direction this points is toward controls that do not require discrimination. Requiring a second channel for consequential actions works whether or not the first message was correctly classified. Delaying irreversible transactions by an interval sufficient for anomalies to surface works without any judgement about the request. Limiting what a single compromised interaction can achieve works regardless of how convincing the interaction was. These are unfashionable controls, they impose friction that product teams resist, and they have the property that their effectiveness does not degrade as the adversary improves.

One structural response remains and deserves more consideration than it usually receives, because it changes the problem rather than improving the odds within it. Much of the fraud described in this chapter depends on the target organisation inferring identity from behaviour, and inference is exactly what synthesis defeats. Where identity can instead be asserted by a party that verified it directly, the synthesis has nothing to attack. Bank-attested account ownership, government-issued credentials presented cryptographically, and employer-attested employment all move the verification burden onto an institution that performed a check in the physical world.

Digital identity infrastructure of this kind has been slow to arrive for reasons that are political and commercial rather than technical, and the fraud economics described here are among the stronger arguments for accelerating it. The costs are also real and should

be stated plainly: attested identity concentrates power in the attesting institutions, creates exclusion for people who lack the relevant relationships, and raises privacy questions about correlation across services that presenting a document does not. A fraud team is not the right body to resolve those questions, and it should understand that the technical direction its incentives point toward carries them.

A closing observation about attribution, since it affects how organisations should reason about who is attacking them. Traditional threat intelligence rests on the assumption that capability is a proxy for resources: a sophisticated intrusion implies a well-funded group, because sophistication required expertise and expertise was scarce. That inference is weakening. A single operator with access to commodity models and a modest budget can now run campaigns whose polish would previously have implied a professional team. Defenders who calibrate their response to the apparent sophistication of an attack risk over-attributing to organised groups and, more damagingly, dismissing the possibility that a highly polished campaign came from someone with no particular resources at all. The next chapter examines a related shift, in which the attacker's objective moves from obtaining money to causing an organisation to spend it.

Chapter 5: Denial of Wallet (DoW) and FinOps Security

Security and financial management have historically been adjacent disciplines with a thin relationship. Security teams considered cost when justifying budget, financial operations teams considered security when a breach produced a line item, and the two rarely met on a technical question. Consumption-based billing for agentic workloads has ended that separation, because the mechanism by which an attacker damages an organisation can now be the organisation's own invoice. This chapter treats that mechanism as a first-class attack class, examines why the engineering properties that make agents reliable also make them financially exploitable, and describes the control layer that has emerged in response.

A word first about why the volume figures in the previous chapter and the cost figures in this one belong in the same argument. An attacker who can generate unlimited convincing interaction at near-zero marginal cost has two ways to profit from that capability. They can use it to extract value, which is fraud, and they can use it to impose cost, which is what this chapter examines. The second has historically been the less attractive option, because imposing cost on a target rarely benefits the attacker directly and was expensive to attempt. Consumption billing has changed the arithmetic: the cost imposed per unit of attacker effort is now large enough that extortion, competitive sabotage, and simple vandalism all become economically coherent.

5.1 The Mechanics of Denial of Wallet (DoW)

Denial of service attacks exhaust a finite resource: connections, memory, bandwidth, worker threads. The defence is capacity and filtering, and the attack ends when the traffic stops. Cloud elasticity changed the shape of this by removing the finite ceiling, which was widely treated as an improvement. Under elastic provisioning a traffic surge is absorbed rather than refused, and the service stays available. The cost of absorbing it appears later, on a bill.

Denial of wallet attacks target that bill directly. The objective is to force the target's own systems into expensive processing paths, so that the organisation pays for work it did not want done. The pattern predates agentic systems, and serverless platforms have been vulnerable to it since function-level billing was introduced, but agentic architectures make it substantially easier to trigger and considerably more expensive when triggered.

The cost of an agentic execution is straightforward to model. For a run of N steps, with T_{in} tokens consumed as input per step and T_{out} tokens generated per step, priced at P_{in} and P_{out} per token respectively, the total cost C is given by $C = N \times (T_{in} \times P_{in} + T_{out} \times P_{out})$. Every term on the right side is under some degree of attacker influence. Input tokens grow when the agent retrieves large documents, which an attacker can arrange by controlling what it retrieves. Output tokens grow when the agent produces long responses, which verbose instructions encourage. The pricing terms rise when the agent escalates to a larger model, which ambiguous material tends to provoke. And the step count N , which multiplies everything else, has no natural upper bound in a system designed to continue working until its objective is met.

An attacker's most efficient move is therefore to attack N . Consider an instruction of the following shape, delivered through any of the injection paths described in Chapter 2: locate and reconcile every matching receipt in the archive, and where any record is missing, generate alternative search parameters and continue until all records are found. The instruction is grammatical, sounds like a legitimate business request, and contains a termination condition that can never be satisfied. An agent given this objective will search, fail to find everything, generate new parameters, search again, and continue. Each iteration is a well-formed request from an authorised client. The loop terminates when a human notices or when the budget is exhausted.

Standard protective controls do not engage against this pattern, and the reasons are worth being precise about because organisations often assume they are covered. Rate limiting throttles requests per unit time, which slows the burn without stopping it, and

rate limits sized for legitimate agentic workloads are necessarily generous because legitimate agents make many calls. Authentication passes, because the calls carry valid credentials issued to an authorised orchestrator. Anomaly detection based on volume may fire eventually, though the threshold has to be set above normal agentic activity, which is already high. Budget alerts fire after the money is spent, and alert delivery in most organisations is measured in tens of minutes while token consumption is measured in seconds.

The severity of the exposure depends on a configuration detail that many deployments have never examined: whether the agent's loop has a hard step ceiling. A deployment with a maximum iteration count has bounded its worst case at the cost of occasionally failing a legitimate long task. A deployment without one has an unbounded worst case. In reviewing agentic deployments, this single question separates organisations that face a contained incident from organisations that face an open-ended one, and it can be answered in a few minutes by reading a configuration file.

Retrieval-augmented architectures introduce a further amplifier that deserves specific mention, because it multiplies input tokens rather than step counts and therefore escapes both of the obvious ceilings. Each step of a retrieval-augmented agent begins by pulling context from a knowledge base, and the volume pulled is determined by a configured retrieval depth rather than by the size of the question. An agent configured to retrieve twenty documents per step consumes twenty documents worth of input tokens on every iteration of a loop, which means a loop that would have been affordable becomes expensive by a factor of the retrieval depth. Reviewing retrieval configuration for interaction with loop behaviour takes a few minutes and is rarely done.

Multi-agent deployments amplify differently and more severely, because the cost of a runaway loop scales with the number of agents participating in it. A planner that decomposes a failing task across sub-agents, each of which encounters its own failure and begins its own recovery sequence, produces a branching consumption pattern rather than a linear one. Practitioners who have observed this describe cost curves that

rise faster than any budget alert threshold was configured to catch, because the alerting assumed linear growth. Where budget ceilings are applied per agent rather than per originating task, the ceilings are individually respected and collectively meaningless.

There is also a subtler variant that targets the pricing terms rather than the step count. An attacker who can influence what an agent retrieves can arrange for it to retrieve very large documents, or documents whose ambiguity provokes escalation to a more expensive model. Neither produces a loop, so step-ceiling protection does nothing, and the cost per step rises by an order of magnitude while the agent behaves entirely normally. Defending against this requires metering at the token level rather than the request level, which is the distinction the control layer described in section 5.3 exists to enforce.

Three properties of the exposure make it awkward to reason about and are worth stating separately, since organisations tend to recognise one and miss the others. The first is asymmetry of effort. The attacker supplies a single instruction and spends nothing further; the target's own infrastructure performs all the expensive work. There is no sustained traffic for a defender to filter, no connection to drop, and no source address to block, because after the initial injection the attacker has disengaged entirely.

The second is delayed visibility. Conventional attacks produce their effect at the moment of execution, and a defender watching the right dashboard sees it happen. Consumption damage accrues silently and becomes visible when billing data is reconciled, which in most organisations happens daily at best and monthly at worst. An attack that begins on a Friday evening may not be observed until the following week, by which point the loop has run for days.

The third is the difficulty of establishing intent after the fact. When a bill arrives showing an anomalous spike, distinguishing a deliberate attack from a badly written prompt from a genuinely difficult task requires the execution trace, and the execution trace is exactly what most deployments discard. Organisations that lack it end up

treating every cost anomaly as an engineering defect, which means deliberate attacks are quietly filed as bugs and the pattern is never recognised.

Ownership compounds the problem organisationally. Cost anomalies are routed to the financial operations function or to the platform team that owns the budget, and those teams investigate them as capacity or efficiency questions. Security is generally not in the notification path for a spending alert, and there is no reason it would be, because spending alerts have never been security telemetry before. Establishing that path, so that a cost anomaly above some threshold generates a security notification alongside the financial one, costs almost nothing and is among the higher-value changes available to an organisation running agentic workloads.

5.2 The Resilience Paradox

Engineers build agents to recover from failure, and they are right to. A system that halts on the first malformed input, transient timeout, or unexpected schema is a system that requires constant human attention, which removes the reason for deploying it. Retry logic, fallback strategies, alternative tool selection, and reformulation of failed queries are all deliberate design choices that make an agent useful in an environment where dependencies are imperfect.

Under consumption pricing, every one of those choices is also a cost multiplier, and their combination produces a failure mode that has no analogue in conventional systems. A conventional program encountering an error either handles it cheaply or terminates, and both outcomes are inexpensive. An agent encountering an error reasons about it, and reasoning costs money. When the reasoning produces a hypothesis that is also wrong, the agent reasons about that, and so on, with each layer of recovery consuming more tokens than the layer beneath it.

Each step is locally reasonable; the sequence has no terminating condition

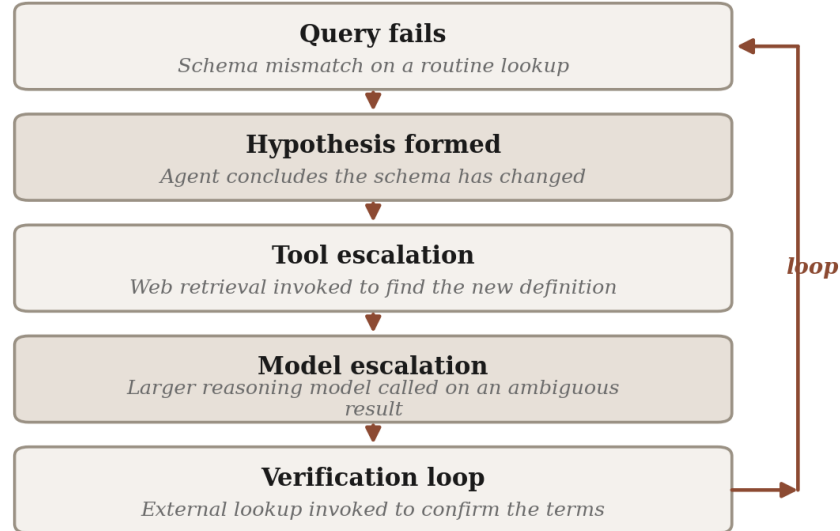


Figure 5.1 The resilience trap. Recovery behaviour that is individually sensible compounds into an unbounded sequence.

Follow the sequence in Figure 5.1 and observe that no step is a mistake. The agent encountering a schema error concludes the schema may have changed, which is a reasonable inference. Deciding to look up current documentation is a reasonable response to that inference. Escalating to a more capable model when the documentation proves ambiguous is exactly what an engineer would want. Consulting an authoritative source to verify an ambiguous term is good practice. Each decision is defensible, the sequence has no termination condition, and the cost per step increases as the agent reaches for progressively more expensive capability.

Financial analysis of enterprise cloud spending during 2026 characterised the aggregate effect as roughly a four hundred million dollar leak in unbudgeted cloud consumption across the Fortune 500, attributed largely to agentic execution that ran longer and consumed more than the deploying teams anticipated. The figure is an estimate assembled from spending patterns rather than a measured total, and the estimation method deserves scrutiny before the number is quoted in a board paper. What is not in dispute among practitioners is the direction and the mechanism: agentic workloads

consume unpredictably, the unpredictability is concentrated in error paths, and error paths are the least tested part of most deployments.

The overnight case is the one that produces the largest individual incidents. An agent scheduled to run a batch process outside business hours encounters an error at midnight, begins its recovery sequence, and continues through it for seven hours with nobody watching. Organisations that have experienced this describe arriving to find a bill that exceeded the month's forecast, generated by a task that was never completed. No attacker was involved, which is the point worth taking from it: the mechanism operates on its own, and an attacker who wants to trigger it deliberately need only supply the initial error.

Quantifying an organisation's own exposure is a short exercise and produces a number that focuses attention better than any general argument. Take the most expensive model in use, multiply its per-token output price by a plausible tokens-per-step figure for the deployment, multiply by the number of steps the agent could execute in an hour at observed throughput, and multiply again by the number of agents that could plausibly be triggered by the same injected instruction. The result is the hourly burn of an uncontained incident. Teams performing this calculation for the first time frequently discover that the figure exceeds their monthly platform budget, and the discovery tends to resolve arguments about whether a mediation layer is worth building.

The same calculation identifies where ceilings should be placed, which is a more useful output than the headline number. If the exposure is dominated by step count, a step ceiling addresses most of it. If it is dominated by per-step cost because the deployment escalates to expensive models on ambiguity, then constraining model selection matters more than constraining iteration. If it is dominated by fan-out across agents, then the ceiling has to be applied at the originating task rather than at the individual agent, which is a different piece of engineering. Deployments differ enough that the generic advice to set a budget limit is close to useless without this analysis.

This is where security and financial operations become the same discipline rather than adjacent ones. The telemetry that would detect a denial of wallet attack, meaning token consumption per task, step counts per session, and the semantic similarity of consecutive steps, is identical to the telemetry that a financial operations team needs for cost attribution. The control that would stop an attack, meaning a budget ceiling per task with a hard stop, is identical to the control that prevents accidental overspend. Organisations that have built this capability generally built it for cost reasons and discovered the security benefit afterwards, which suggests a practical route for security architects who find the security case hard to fund.

A related design tension appears when teams attempt to tune the recovery behaviour itself. The obvious response to the resilience trap is to make agents give up sooner, and the obvious problem with that response is that it makes them less useful in exactly the situations where autonomy was supposed to help. An agent that abandons a task on the first transient failure requires the human supervision the deployment was meant to remove. The tuning question has no general answer, and the pattern that has worked in practice is to make persistence conditional on progress rather than on an attempt count, permitting many attempts while the agent's state keeps changing and stopping quickly once it starts repeating itself.

Measuring progress is harder than counting attempts, which is why this pattern is less common than it should be. A workable approximation compares the semantic content of successive steps: an agent that is genuinely working through a problem produces steps that differ from one another, while an agent that is stuck produces variations on a theme. This is the same measurement that section 5.3 describes as a detection control, and using it as a termination condition inside the agent rather than only as an external check turns a monitoring signal into a design property. The distinction matters because an external control stops a runaway after it starts, whereas a progress condition prevents most of them from starting.

5.3 Mitigation: Cloud Cost Arbitrage Layers

The control that has emerged is a mediation layer positioned between the orchestration engine and the model provider interfaces. Products in this category, including Mosaic Sentinel, operate as proxies through which all model calls pass, and they exist because the enforcement point has to sit somewhere that sees every call and can refuse one. The orchestrator cannot police itself, since a compromised or looping orchestrator is the thing being policed, and the provider will not, since the provider is being paid.

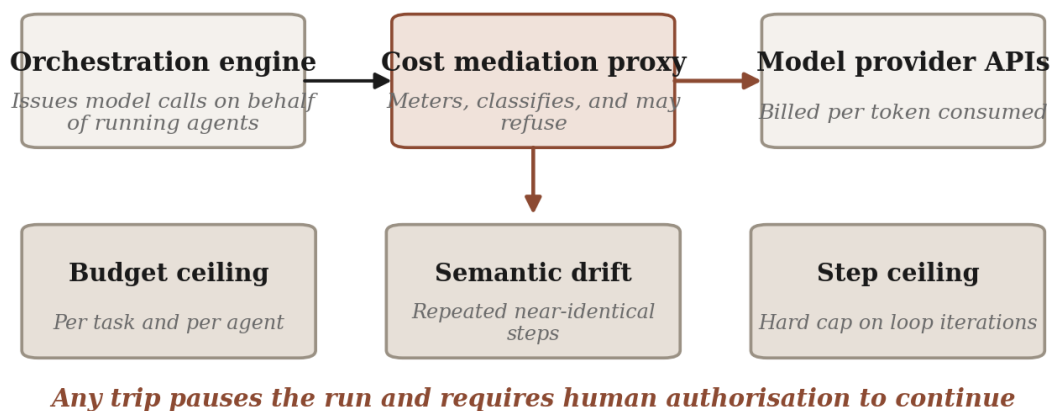


Figure 5.2 A cost mediation proxy positioned between the orchestrator and provider interfaces, applying three independent ceilings.

Three enforcement mechanisms operate at this layer, and they are independent, which matters because each catches a case the others miss.

Budget ceilings are the simplest. A task, an agent, or a team is allocated a spending limit, consumption is metered in real time against it, and the proxy refuses further calls when the limit is reached. A representative configuration might cap a single task at fifty dollars, which is generous for legitimate work and stops a runaway loop within a few minutes. The design question is what happens at the ceiling, and the answer that works is a pause requiring human authorisation rather than a hard failure, since legitimate tasks do occasionally need more and a control that fails destructively will be disabled by the first team it inconveniences.

Step ceilings bound the iteration count directly. This catches the case where an agent loops cheaply, consuming little per step and a great deal in aggregate, which a budget ceiling would eventually catch and a step ceiling catches sooner. Step ceilings are also easier to reason about during design, because a team can usually say how many steps a task should reasonably require, whereas estimating its cost in advance is harder.

Semantic drift analysis is the most interesting of the three and the least widely deployed. The proxy compares consecutive steps for similarity, measuring whether the agent is making progress or repeating itself with cosmetic variation. A healthy execution moves through distinguishable states. A trapped execution produces steps that resemble each other closely, and that resemblance is measurable using the same embedding machinery the agent already uses. Repetition above a threshold terminates the run. The mechanism catches loops before either ceiling is approached, and it produces a signal that is diagnostically useful afterwards, since the point at which similarity spiked identifies where the agent became confused.

Two implementation cautions apply. The first is that a mediation proxy is now a component on the critical path of every model call, with the availability and latency obligations that position implies. Organisations have deployed these layers and then found them contributing more downtime than the attacks they prevent, which is a poor trade. The second is that the proxy sees every prompt and every response, making it among the most sensitive data flows in the environment, and its own access controls and retention policy deserve the scrutiny given to a logging system holding production data.

Two further controls belong in this layer and are frequently omitted because they address correctness rather than cost. The first is a wall-clock ceiling on task duration, independent of steps and spend, which catches the case of an agent blocked on a slow external dependency and retrying patiently for hours at low cost per attempt. The second is a cap on the total volume of content an agent may retrieve within a session, which bounds the input-token amplification described in section 5.1 and has the useful

side effect of surfacing retrieval misconfigurations that would otherwise show up only on the bill.

Alerting design deserves a moment as well, because the natural configuration is wrong in a specific way. Thresholds set as a percentage of budget assume that consumption is roughly linear and that crossing eighty percent leaves time to react. Agentic consumption is bursty, and the interval between eighty percent and two hundred percent can be shorter than the alert delivery path. Effective alerting for this workload triggers on rate of change rather than on level, so that a sudden acceleration generates a notification while the absolute number is still small, and it delivers to a channel that is monitored continuously rather than to an email address someone reads in the morning.

Contractual and procurement questions belong in this chapter as well, since a proportion of the exposure can be moved rather than engineered away. Model providers offer spending caps, prepaid credit arrangements, and hard account limits with varying degrees of granularity, and the difference between a provider that enforces a per-key ceiling in real time and one that reconciles overage monthly is the difference between a contained incident and an open-ended one. Procurement teams negotiating these agreements are rarely briefed to ask, and the question belongs on the security review checklist for any model provider relationship.

The insurance position is less settled and is moving. Cyber policies were drafted against a threat model in which loss arises from unavailability, data disclosure, or extortion, and a claim for excess cloud consumption caused by a maliciously induced processing loop does not fit those categories cleanly. Organisations carrying material exposure of this kind should raise it with their broker before an incident rather than discovering the coverage position during one, and should expect the answer to be unsatisfying for another cycle or two while underwriters develop a view.

A final note on how this class will develop. Every mitigation in this chapter constrains consumption, and each constraint is a place where an attacker can instead aim for denial of service by triggering the ceiling deliberately. An agent halted because its

budget cap tripped is an agent that has stopped working, and an attacker who can trip the cap cheaply has converted a financial attack into an availability one. Designing the ceilings so that they contain a single task rather than halting a shared budget, and so that the pause requires authorisation rather than terminating the workload, is what keeps the remedy from becoming the next vulnerability.

The deeper mitigation is architectural rather than defensive. An organisation whose agentic workload runs against open-weight models on its own infrastructure has a fixed cost base, and a runaway loop consumes capacity that was already paid for rather than generating an unbounded invoice. The failure becomes an availability incident, which security teams have thirty years of practice handling, instead of a financial one. This trade is examined properly in Chapter 13, where the pricing structures and their security implications are treated together. The next chapter returns to the technical surface and examines the communication paths between agents, where the trust assumptions inherited from service architecture turn out to be the weakest link in most multi-agent deployments.

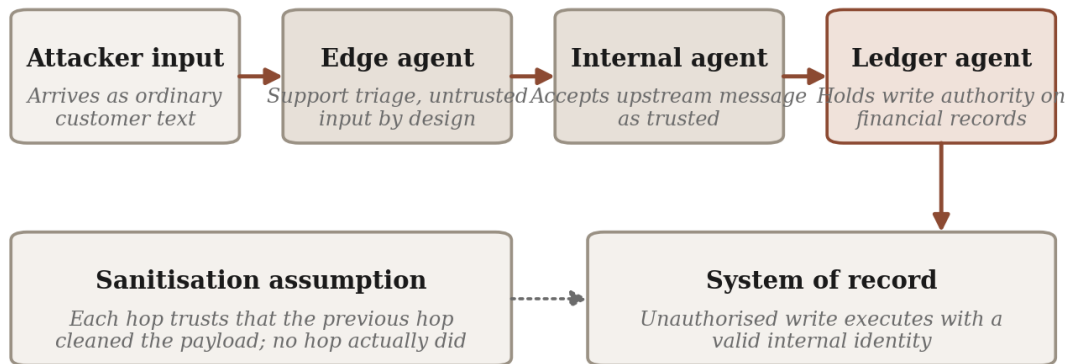
Chapter 6: Secure Agent-to-Agent (A2A) Communications

Single-agent deployments are the exception in serious production use. Complex objectives are decomposed across specialised components: a planner that breaks down the work, executors that perform parts of it, retrieval agents that gather context, review agents that check output before it is released. The design is sound, mirrors how organisations divide work among people, and produces better results than asking one agent to do everything. It also creates a communication surface whose security properties have received considerably less attention than the surface at the system's edge, and that surface is where the most serious multi-agent failures have occurred.

6.1 The Risk of Ungoverned Peer-to-Peer Talk

Service-oriented architecture supplied the trust assumption that causes the problem from service-oriented architecture, where it was correct. In a conventional microservice environment, a service receiving a request from another internal service can reasonably treat the input as validated, because the calling service was written by the same organisation, its behaviour is deterministic, and if it passes malformed data that is a bug with an owner and a fix. Defensive validation at every internal boundary is possible and is sometimes done, and it is usually judged an inefficient use of engineering effort relative to the risk.

Multi-agent systems inherit the assumption and invalidate its basis. An agent receiving a message from another agent is receiving text that the sending agent generated, and the sending agent generated it after processing whatever content it had been exposed to. If the sending agent read a customer ticket containing an injected instruction, that instruction can propagate into the message it sends downstream, either verbatim or in paraphrase. The receiving agent, applying the inherited assumption that internal traffic is clean, processes it as a trusted instruction from a peer.



Trust is inherited along the chain while the payload travels unchanged

Figure 6.1 Cascading propagation. Each hop assumes an earlier hop performed sanitisation, and the assumption is never checked.

Propagation dynamics make this worse than the single-agent case in a specific way. Agents are typically arranged so that edge-facing components have narrow permissions and internal components have broad ones, which is correct design. A support triage agent that reads customer messages holds almost no authority. A ledger agent that writes financial records holds a great deal. The edge agent is exposed and harmless; the internal agent is powerful and, by assumption, protected by its position. Cascading injection defeats exactly this arrangement, because the payload travels from the exposed component to the powerful one along a path the architecture provided, carried in a message the receiving agent has every reason to trust.

Researchers at Unit 42 and Lares Labs have documented this class under the name A2A session smuggling, describing exploits in which a compromised or malicious sub-agent manipulates the task graph of a trusted orchestrator. The mechanism generalises beyond the specific cases published. Any arrangement in which one agent's output becomes another agent's input, without validation at the boundary, permits an instruction introduced at any point to reach any component downstream of it.

Adoption of the Model Context Protocol has made this both more common and more consequential, and the reasons are structural rather than a criticism of the specification. A common connector standard lowered the cost of wiring agents to tools and to each other, which is the point of a standard and has produced real benefit. Lower integration cost produces denser connection graphs, and a denser graph means more paths along which a payload can travel. The protocol as specified addresses connection and message format rather than authorisation, and it carries no native mechanism for constraining what an authenticated caller may request once connected. Organisations that adopted it as an integration convenience frequently did not observe that they had also adopted a lateral movement surface.

Three properties make this surface difficult to defend with existing tooling. The first is that inter-agent traffic is east-west and internal, so perimeter controls do not see it, for the reasons set out in Chapter 1. The second is that the traffic is semantically dense: an agent-to-agent message is natural language carrying an instruction, a context, and often a delegation of authority, and no schema validation can determine whether its content is legitimate. The third is that the topology is dynamic. Agents create sub-agents at runtime, so the communication graph at any moment reflects decisions a planner made seconds earlier, and a network policy written against a static topology will not describe the system as it actually runs.

Visibility into this surface is correspondingly poor. Survey work covered in Chapter 8 found that under a quarter of organisations maintain active visibility into the pathways their agents use to communicate with each other. The remainder operate with what practitioners have taken to calling dark agent traffic: workflows that exchange context, delegate tasks, and pass credentials entirely outside the monitoring scope of the security operations function. An attacker who compromises any component in such an environment can move through the mesh without generating a single record that anyone will read.

Detection on this surface is harder than on the external one for a reason that is easy to state and difficult to engineer around. At the system edge, a defender knows which side of the boundary is untrusted, and can apply asymmetric scrutiny accordingly. Inside a multi-agent mesh, every component is simultaneously a producer and a consumer of untrusted content, and there is no side to point the scrutiny at. Applying full evaluation to every internal message multiplies the deployment's inference cost by the density of its communication graph, which for a system doing genuine multi-agent work is a large multiplier.

This drives the tiering argument that Chapter 8 develops in detail. Messages that carry a delegation of authority, that reach an agent with irreversible tool access, or that cross from a component exposed to external content into one that is not, warrant evaluation. Messages passing between two agents that both hold narrow read permissions generally do not. Drawing that boundary requires knowing which agents hold which authority, which is the identity work of Chapter 3 appearing again as a prerequisite, and it is the third time in this book that a control has turned out to depend on attribution being solved first.

There is a design principle available here, and it is one that distributed systems engineers will recognise from a different context. In a system where any component may be compromised, no component should trust the assertions of another about authority, and every boundary should validate independently rather than relying on a claim that validation occurred upstream. This is the same reasoning that produced end-to-end argument in network protocol design and mutual authentication in service meshes. Applying it to agents means that a receiving agent treats an inbound message as untrusted input regardless of its origin, validates the authority it claims against a source the sender does not control, and constrains its own actions by its own scope rather than by what the message requests.

A concrete example clarifies what independent validation means at this boundary, since the principle is easy to state and its implementation is frequently misunderstood.

Suppose a retrieval agent returns a document summary to a planning agent, and the summary includes a sentence stating that the finance agent has been authorised to approve payments above the usual threshold for this vendor. Content validation asks whether that sentence looks like an injected instruction, which is a judgement call that will sometimes be wrong. Authority validation asks a different question: does the finance agent's current credential permit approving payments above the threshold. That question has a definite answer, it is answered by the identity infrastructure rather than by a model, and the answer does not depend on how convincingly the sentence was phrased.

A general form of this move appears throughout the book and is worth naming explicitly here because Part III depends on it. Wherever a security decision can be reframed from a question about the content of a message to a question about the authority of an actor, it should be, because the first question is answered by a component that can be deceived and the second by a component that cannot. Content evaluation retains a role in detection, where its errors are tolerable because a false negative means a missed alert rather than an executed transaction. Placing content evaluation in the enforcement path, where a false negative means the action proceeds, puts a probabilistic component in a position that requires a deterministic one.

Implementing that principle costs something, and the cost falls on capability rather than on infrastructure. Multi-agent systems derive much of their value from flexible delegation, meaning an agent that encounters a task outside its competence hands it to one that can perform it. Constraining delegation to pre-declared paths with validated authority makes the system more predictable and less adaptive. Most organisations that have worked through this have concluded that the constraint is acceptable for the small set of agents holding consequential authority and unnecessary for the rest, which is the same reversibility-based boundary that has appeared in each preceding chapter.

Two further attack shapes on this surface deserve description, because they are less obvious than straightforward payload propagation and neither is addressed by

validating message content alone. The first is authority laundering. An edge agent holding narrow permissions cannot perform a consequential action itself, and it can request one. If the request is phrased as a routine delegation and the receiving agent has been designed to accept work from peers, the low-privilege component has obtained the effect of a high-privilege action without ever holding the privilege. Every credential check in the path passes, because the powerful action was performed by the agent authorised to perform it, on the basis of a request from a component the architecture told it to trust.

The second is orchestrator manipulation. In arrangements where a planning agent decides which sub-agents to invoke and in what order, the planner's decisions are made from context that includes results returned by those sub-agents. A compromised sub-agent that returns crafted results can therefore influence which agents the planner calls next and what it asks them to do. The compromise moves upward through the hierarchy, from a leaf component to the coordinator, which is the opposite of the direction most threat models assume and is the mechanism the session smuggling research describes.

Defending both requires the receiving side to evaluate requests against its own understanding of the current task rather than against the requester's claim about it. Practically, that means an inbound delegation carries a reference to the originating session and its declared scope, and the receiver checks the requested action against that scope rather than against a general trust relationship with the sender. This is more than message validation and less than full semantic evaluation: it is a scope check, it is cheap, and it closes the laundering path because a request outside the originating session's scope is refused regardless of which internal component made it.

Topology discipline provides a second layer. Multi-agent architectures assembled organically tend toward dense connection graphs in which most agents can reach most others, because each connection was added to solve a problem and none were removed. Declaring the permitted communication paths explicitly, so that an agent may only

receive delegations from a named set of peers, converts the graph from whatever it happens to be into something that can be reviewed. Most deployments discover on performing this exercise that a substantial fraction of their existing paths are unused, and removing them is free risk reduction.

6.2 Shifting from Shared API Keys to Dynamic A2A Identity

Behind the trust problem sits a credential problem, and the survey data on it is stark. According to the Gravitee 2026 report, 45.6 percent of organisations still use shared static API keys for communication between agents, while only 21.9 percent treat agents as independent identity-bearing entities. Slightly under half of the surveyed population has, in effect, given its agents a single shared password.

Consequences of the shared key model compound one another. Attribution fails first: when every agent presents the same credential, no log entry identifies which agent made a call, and an investigation cannot establish where a compromise entered. Scope fails next: a shared credential must carry the union of all permissions any holder requires, so the least privileged agent in the group holds the permissions of the most privileged. Revocation fails third and most visibly: rotating a compromised shared key requires updating every agent that uses it simultaneously, which in practice means an outage, which in practice means the rotation is deferred. And lateral movement becomes free, since an attacker who extracts the key from any single agent holds the authority of all of them.

Each of these is a well-understood weakness that security teams eliminated from human access management two decades ago and from service architecture more recently. Their reappearance in agentic deployments reflects deployment speed rather than ignorance. A shared key is what a team reaches for when they are proving a concept, and the concept succeeded and went to production without anyone revisiting the decision.

Property	Shared static keys (45.6 percent)	Dynamic session-bound identity (21.9 percent)
Credential lifetime	Permanent until manually rotated	Minutes, bound to a single session
Access scope	Union of every holder requirement	Narrowed to the delegated task
Attribution	Calling agent cannot be identified	Verifiable chain naming agent and principal
Revocation	Requires coordinated rotation across holders	Expiry is automatic; renewal can be refused
Lateral movement	Unrestricted once a key is obtained	Bounded by scope and by expiry
Audit position	Records show that something authorised acted	Records show which component acted and why

Table 6.1 Credential models for agent-to-agent communication and their operational properties.

A better arrangement extends the delegation architecture of Chapter 3 across the internal boundary. Each agent authenticates individually using a short-lived credential bound to its workload identity and to the current session. When one agent delegates to another, it does not pass its own credential. It requests a new, narrower credential for the receiving agent, scoped to the delegated subtask and carrying the delegation chain that records how the authority was derived. The receiving agent presents this credential to whatever resources it needs, and those resources validate the chain independently.

Two properties of this construction are worth separating because implementations often achieve one without the other. Attenuation means that authority can only narrow as it propagates: an agent cannot delegate more than it holds, and each hop reduces scope. This bounds cascading compromise structurally, since a payload that reaches the fourth agent in a chain arrives with the authority of that agent's narrow subtask rather than the authority of the orchestrator. Verifiability means that the resource server can evaluate the entire chain rather than trusting the immediate caller's assertion, which removes the propagated-trust failure that Figure 6.1 illustrates.

Message-level integrity belongs alongside credential-level identity and is frequently overlooked, since teams who have solved authentication assume the message content is thereby protected. It is not. An authenticated channel confirms who sent a message and

prevents a third party from altering it in transit; it says nothing about whether the message content faithfully represents the task it claims to convey. Signing the delegation payload itself, so that the subtask, its scope, and its originating session are covered by a signature the receiving agent verifies, closes the gap between a trusted channel and a trustworthy instruction.

Replay deserves a similar mention. A delegation message captured from a legitimate exchange and resubmitted later can trigger the same action again, and in a financial context the same action twice is an incident. Standard countermeasures apply without modification: a nonce, a timestamp with a narrow acceptance window, and a record of recently processed message identifiers. These are unremarkable requirements for any authenticated protocol, and their absence from a number of early multi-agent implementations reflects how quickly those implementations were assembled rather than any genuine difficulty.

Salt Security and other vendors in the API security space have built products around this model, and the underlying mechanisms are open standards rather than proprietary inventions. A team can construct the architecture with an identity provider supporting token exchange, a workload attestation service, and a gateway capable of validating chains. The engineering effort is real, and it is integration work rather than research.

Where a full delegation architecture is out of reach in the near term, an intermediate position captures a meaningful share of the benefit. Giving each agent a distinct credential, even a long-lived one, restores attribution and permits per-agent scoping, and those two properties together eliminate the worst consequences of the shared key model. Short lifetimes and chain validation can follow. The sequencing argument from Chapter 3 applies unchanged: distinct identities first, narrow scopes second, short lifetimes third, full attestation last, with value delivered at each stage rather than only at completion.

Operational objections to short-lived inter-agent credentials are usually raised in one of two forms, and both have reasonable answers. The first concerns long-running tasks: an

agent working on an objective that takes hours cannot hold a credential that expires in minutes. The answer is that it holds a succession of them, renewing as it proceeds, which is how every other long-running authenticated workload in a modern environment already behaves. The renewal path becomes a control point, since an orchestrator that has been marked as compromised simply fails to renew.

The second concerns availability, and it carries more weight. Making every inter-agent message depend on a credential issuance service means that service becomes a dependency of the entire agentic estate, and its failure stops everything. The mitigation is the standard one for identity infrastructure, meaning redundancy, regional distribution, and cached credentials with bounded lifetime that permit continued operation through a short outage. What organisations should avoid is the tempting middle position of falling back to a shared static key when issuance is unavailable, because that fallback is permanently present in the configuration and an attacker who can disrupt the identity service has obtained it.

A final observation about where this part has arrived. Chapters 4 through 6 have described three ways in which autonomy benefits an attacker: it removes the trade-off between scale and precision in fraud and social engineering, it converts a target's own resilience mechanisms into a financial weapon, and it creates internal communication paths along which a single injection can reach components it should never touch. What unites them is that in each case the attacker's activity, viewed by any individual control, appears legitimate. The fraudulent profile behaves like a customer, the runaway loop behaves like a difficult task, and the cascading payload behaves like an internal message.

A short review procedure follows from everything in this chapter and can be run against an existing deployment in an afternoon. Enumerate the agents. For each, record which other agents may send it work, what credential it presents to downstream resources, and whether that credential is distinct to it. Mark every agent that holds an irreversible capability, using the reversibility test from Chapter 2. Then trace, for each marked agent, every path by which externally originated text could reach it, counting hops

through other agents as paths. The output is a short list of the routes along which a single injected instruction could reach a consequential action.

Most organisations running this exercise find between two and five such routes, and find that at least one of them was unknown to the people who built the system, because it runs through a delegation added later for an unrelated reason. Closing those routes, by validating scope at the receiving end or by removing the delegation entirely, addresses the specific exposure that the whole of this chapter describes, and it requires no new products. The exercise is worth repeating whenever the agent population changes materially, which in an active deployment means quarterly at least.

A note on enforcement placement before this part closes. The credential work described here establishes who is calling and with what authority, which is necessary and insufficient. A validly issued, correctly scoped, cryptographically attested credential says nothing about whether the request it accompanies serves the task it was issued for, and every attack described in this part operates through requests that are valid in exactly that sense. The remaining question is how an enforcement layer decides whether an authorised action is a legitimate one, which requires observing behaviour rather than checking credentials. Part III takes up that question, beginning with why the detection methods enterprises already own are aimed at the wrong properties.

Part Three: Detecting the Ephemeral

Chapter 7: Rethinking Anomaly Detection

Parts I and II established that agentic compromise produces no syntactic signature: the credentials are valid, the protocols are correct, the destinations are permitted, and the actions are within the technical authority of the acting component. Everything that distinguishes an incident from ordinary work is a relationship between the action and its purpose. This part addresses the practical question that follows, which is how an organisation detects something whose only distinguishing property is meaning.

Detection is where security teams have historically spent most of their operational budget, and the tooling in place represents years of tuning against a threat model that agentic systems break. This chapter explains precisely how the existing model fails, which parts of it can be salvaged, and what replaces the rest. The answer proposed is transition-state modelling, and the argument for it is that behaviour, unlike identity or content, cannot be forged by an attacker who has hijacked a legitimate component, because the attacker is constrained to act through that component's actual capabilities and every deviation leaves a trace in the sequence.

7.1 The Failure of User-Centric Baselining

User and entity behaviour analytics has been among the more successful ideas in enterprise security over the past decade. The premise is sound: rather than enumerating bad behaviour, which is unbounded, model normal behaviour for each actor and alert on departures from it. An employee who logs in at nine, works from two known locations, touches a stable set of systems, and downloads modest volumes generates a tight distribution, and a departure from that distribution is informative. The technique catches compromised accounts that signature-based tools miss, because the credential is valid and only the pattern of use is wrong.

Every variable that makes this work is a property of human beings, and none of them survives contact with an autonomous workload.

Time of activity carries information because people sleep. An employee authenticating at three in the morning is either working unusually or not the employee. Agents have no diurnal pattern by design, since running continuously is among the reasons they were deployed, and an agent active at three in the morning is an agent doing its job. The variable does not merely become less informative here; it becomes empty, because there is no time at which activity would be unexpected.

Location carries information because people occupy one place at a time and travel at bounded speed. The impossible-travel detection built on this reasoning is among the most reliable signals in enterprise security. Agents execute across elastic infrastructure, and the same logical workload may present from several regions within a minute as instances are scheduled, which produces an impossible-travel alert on every ordinary morning. Teams that have connected agent identities to existing analytics platforms describe exactly this outcome: a flood of alerts that are all false, followed by suppression rules that switch the detection off for the population that most needs monitoring.

Volume carries information because human throughput is bounded. A person who reads four hundred customer records in an hour is doing something other than serving customers. An agent reads four hundred records because that is what the task required, and its baseline volume sits several orders of magnitude above any human threshold. Setting a volume threshold for agents is possible in principle and difficult in practice, because legitimate agentic volume is itself highly variable: the same agent may read two records on one task and forty thousand on the next, and both are correct.

Peer group comparison, which analytics platforms use to compensate for thin individual histories, fails for a subtler reason. The technique works by comparing an actor against others in the same role, on the reasoning that ten accounts payable clerks should behave similarly and the one that does not warrants attention. Applying this to agents requires a notion of role, and agent roles are defined by system prompts that individual teams write without coordination. Two agents nominally performing document review may have entirely different tool sets, retrieval depths, and escalation behaviour, so the peer

group is not a peer group. Where organisations have imposed a shared agent template, peer comparison becomes usable again, which is an argument for platform standardisation on grounds that have nothing to do with maintainability.

Alert triage suffers a related degradation that shows up in operational metrics before anyone diagnoses the cause. Analytics platforms score events and route the highest scores to analysts, and the scoring assumes a population where anomalies are rare. Introducing a population whose ordinary behaviour scores as anomalous inverts that assumption: the queue fills with agent activity, analysts learn within a fortnight that agent alerts are noise, and the learned response transfers to the one agent alert that mattered. This is alert fatigue arriving through a new door, and the damage it does is to the detection of everything else, not only to the detection of agents.

Telemetry footprint: human operator against autonomous agent

Human user		Autonomous agent
	Activity window	
Roughly eight hours		Continuous, no diurnal pattern
	Request pacing	
Episodic bursts with pauses		Millisecond intervals, sustained
	Network origin	
A small set of known addresses		Elastic, multi-region, ephemeral
	Authentication	
Multi-factor, device-bound		Programmatic tokens, no challenge
	Session length	
Minutes to hours		Indefinite, spanning restarts
	Error behaviour	
Stops and asks for help		Retries and reformulates silently

Figure 7.1 Comparison of the telemetry each population generates. Every variable on which conventional baselining depends behaves differently.

Authentication patterns carry information because humans authenticate interactively and interactive authentication generates rich signals: a device, a factor, a challenge response, a session establishment event. Agents present programmatic credentials with no challenge and, where the delegation architecture of Chapter 3 has been implemented,

obtain fresh ones on a schedule. The authentication event that anchors human session analysis has no informative equivalent.

The mixed-population problem compounds all of the above and is the reason many organisations discover the failure only after deployment. Where agents authenticate using credentials issued to people, which the survey data from Chapter 3 suggests is the majority case, the analytics platform computes a baseline across a mixture of human and machine activity attributed to the same identity. The resulting distribution describes neither population. It is too wide to catch a compromised human account, because the machine activity has inflated the normal range, and too narrow to accommodate legitimate agentic variation. Organisations in this position hold detection coverage that is measurably worse than it was before the agents arrived, while their coverage reporting shows no change at all.

Error behaviour is the last of the conventional variables and the one whose loss is least appreciated. A person encountering an unexpected refusal stops and asks a colleague, and the resulting support ticket or help request is itself a security-relevant event that many organisations correlate. An agent encountering the same refusal reformulates and tries a different approach, silently, and continues doing so until it succeeds or exhausts its budget. The visible trace of friction, which security teams have relied on more than they realise, disappears. An attacker probing an agent's boundaries generates no complaints, no tickets, and no puzzled colleagues, and the probing looks identical to an agent doing difficult work.

Something is salvageable from the existing investment, and it is worth identifying before moving on. The correlation infrastructure retains its value: the pipelines that gather events from disparate sources, the storage that retains them, the query capability that lets an analyst reconstruct a sequence. What has to be replaced is the modelling layer sitting on top, meaning the specific features chosen and the notion of a baseline as a distribution over per-session variables. The organisational implication is more

encouraging than it first appears, since the expensive part of a detection programme is the data plumbing and the plumbing survives.

A further point about baselines deserves attention because it determines how quickly a new detection approach can be brought into service. Human behavioural baselines take weeks to establish and remain stable for months, because a person's role changes slowly. Agent baselines can be established in hours, since an agent performing a repeated task generates a large sample quickly, and they become invalid the moment someone edits a system prompt or adds a tool. The operating tempo of the detection function therefore has to match the deployment tempo of the platform team, which means baseline invalidation has to be triggered by configuration change events rather than discovered through alert noise. An organisation whose agent definitions are under version control has the necessary hook available; one whose agents are configured through a console does not.

One class of conventional signal survives the transition intact and deserves to be identified, since the argument so far has been uniformly negative. Signals that measure the relationship between an actor and a resource, rather than the properties of the actor's session, retain their meaning. An agent accessing a system it has never accessed before is informative whether the agent is software or a person. A first access to a data classification the workload has no business reading is informative. A permission grant occurring outside a change window is informative. These are properties of the pairing rather than of the pattern of use, and they do not depend on any assumption about diurnal rhythm or human throughput.

Building on this class first has a practical advantage for teams that need to show progress quickly. First-access detection requires no behavioural modelling, no baseline period, and no new analytics capability. It requires an inventory of which workloads have touched which resources, which most environments can assemble from existing audit logs, and an alert when the set grows. The signal is coarse and it fires on the

specific event that matters most in an agentic compromise, which is a hijacked agent reaching for something outside its ordinary territory.

7.2 Transition-State Anomaly Detection

If per-session variables carry little information, the alternative is to model the sequence. An agent performing a task moves through a series of states: it receives an objective, retrieves context, calls a tool, evaluates a result, calls another tool, produces output. The identity of each state is uninformative on its own, since every state in the sequence is a legitimate operation the agent is permitted to perform. The transitions between them are where the information lives.

The formal object here is a state machine derived from observed behaviour. For a given agent role, the modelling process records which operations follow which, builds the graph of observed transitions, and assigns each a frequency. A support triage agent retrieves a message, classifies it, looks up an account, drafts a response, and returns it. Those transitions recur across thousands of executions, and the resulting graph is compact, because an agent doing a defined job explores a small region of its theoretical possibility space.

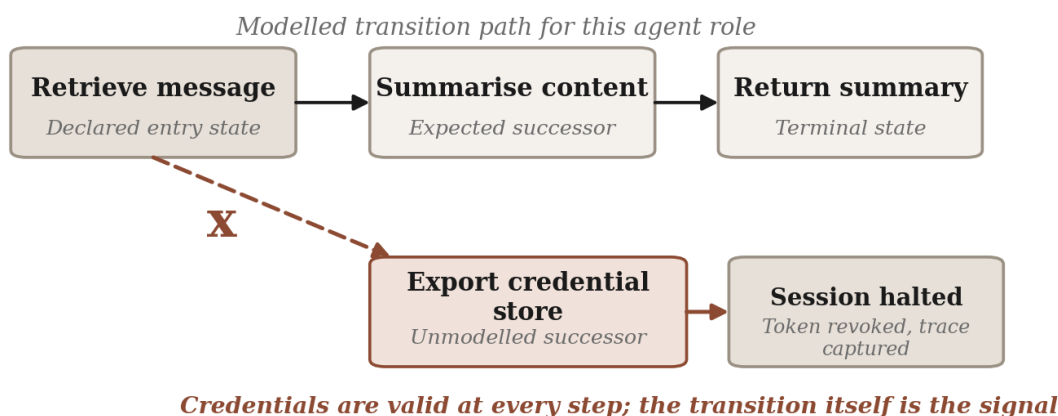


Figure 7.2 Transition-state detection. The anomaly is the edge between states, not either state considered alone.

Consider what this catches that per-request evaluation does not. An agent that retrieves a message and then exports a credential store has performed two operations, each of which it is authorised to perform, in a sequence that has never occurred in its observed history. A request-level control examines the export, confirms the credential permits it, and allows it. A transition-level control examines the edge from message retrieval to credential export, finds no such edge in the model, and refuses. The second control is evaluating something the attacker cannot make look normal, because making it look normal would require the hijacked agent to perform its usual sequence, and performing its usual sequence is not the attack.

Three signal families populate this model, and deployments generally implement them in this order because each is progressively harder to instrument.

Tool-call sequences are the first and simplest. Each agent role has a characteristic vocabulary of operations and a characteristic grammar governing their order. Deviation from the grammar is detectable using techniques borrowed from system-call sequence analysis, a body of work in intrusion detection dating to the 1990s that applies here with little modification. The advantage of starting here is that tool calls are already logged in most orchestration frameworks, so the data collection problem is solved before the modelling begins.

Token consumption rates are the second. Every step consumes a measurable quantity of input and output tokens, and the profile of that consumption across a task is characteristic of the task. A step that consumes ten times its usual input indicates the agent retrieved something unexpectedly large, which is the input amplification pattern from Chapter 5. A sequence of steps with rising consumption indicates escalation, which is the resilience trap. This family has the pleasant property of serving detection and cost management simultaneously, which as Chapter 5 noted often determines whether it gets funded.

Semantic similarity between steps is the third and most informative. Embedding each step's content and measuring the distance between consecutive steps produces a

progress signal. Healthy execution moves through distinguishable states, so consecutive distances are moderate. A trapped agent produces near-identical steps, so distances collapse toward zero. A hijacked agent produces a discontinuity, a single large jump where the task changed, and that jump is the clearest available indicator of the moment an injection took effect. The cost is an embedding computation per step, which is small relative to the inference the step already performed.

The graph itself repays inspection independently of any detection value, which is a benefit teams rarely anticipate. Building the transition model for an agent role produces a picture of what that agent actually does, as distinct from what its designers believed it did, and the two differ more often than not. Teams performing this exercise routinely find agents calling tools nobody remembered granting, retrieving from sources that were meant to be decommissioned, and reaching systems that were assumed to be outside their scope. The model is a discovery instrument before it is a detection control, and the findings from the first construction typically justify the effort on their own.

A useful refinement records the conditions under which each transition occurs, alongside the transition itself. An agent that exports data has a legitimate reason to do so under some circumstances, and the transition into export is unremarkable when preceded by an explicit request and remarkable when it follows a retrieval from an external source. Conditioning the model on the preceding context, rather than treating transitions as memoryless, catches the specific sequence that injection produces and reduces false positives on legitimate variation. The cost is a larger model and more data to train it, which is why deployments generally begin memoryless and add conditioning for the transitions that matter most.

Multi-agent workflows require the model to span components, which is where the approach becomes genuinely harder. A task decomposed across five agents produces five local sequences, each of which may be individually unremarkable while the composition is not. Modelling the composite requires the delegation records described in Chapter 9, so that the analysis can follow a task across the components that handled

it rather than treating each component's activity as an independent stream. Deployments that model per-agent and not per-task will catch a hijacked agent doing something out of character and miss a cascade in which every agent behaves normally for its role.

Storage and computation for this modelling are modest, which is worth stating because teams assume otherwise when a machine learning technique is proposed. A transition graph for an agent role is a small object, typically a few thousand edges, and updating it is an increment. The similarity computation is one embedding per step, using a small model that costs a fraction of what the step itself consumed. The expensive part of this programme is the telemetry collection described in the next chapter, and organisations that have paid for that already possess most of what transition modelling requires.

Three practical difficulties attend this approach, and an honest treatment has to state them rather than presenting the design as settled.

Cold start is the first. A newly deployed agent has no behavioural history, so there is no model to compare against, and the period immediately after deployment is when misconfiguration and unexpected behaviour are most likely. Deployments handle this by beginning in an observation mode where deviations are recorded and not blocked, which leaves a window of reduced protection, or by deriving an initial model from the agent's declared tool set and permitted scope, which is less accurate and available immediately. The second approach connects the detection model to the declarative scoping discussed in Chapter 3, and organisations that implemented scoping find they have most of an initial behavioural model already written down.

Legitimate novelty is the second. An agent asked to do something genuinely new produces a transition the model has not seen, and the correct response is to permit it, since the flexibility to handle novel requests is why an agent was deployed rather than a script. The distinction between an unusual legitimate task and a hijacked one is not always recoverable from the transition alone. What works in practice is to weight the response by the consequence of the action rather than by the novelty of the transition,

which returns to the reversibility criterion that has appeared in every preceding chapter. A novel transition into a reversible action is logged; a novel transition into an irreversible one is held for confirmation.

Model poisoning is the third and the most serious, because it is the failure mode an intelligent adversary will aim for. A behavioural baseline learned from observed activity can be shaped by an attacker patient enough to introduce the target behaviour gradually while the model is learning. This is the memory injection problem of Chapter 2 relocated into the detection layer, and it has the same remedy: separate the authority to influence the model from the authority to act, so that baseline updates require review rather than accruing automatically from observed traffic. Deployments that retrain continuously on live activity have built an attack surface into their detection system, and the convenience of continuous retraining is not worth what it costs.

Evaluating whether this approach works requires a measurement discipline that most detection programmes lack, and the difficulty deserves acknowledgement. Conventional detection efficacy is assessed against a corpus of known attacks, and no such corpus exists for agentic compromise at any useful scale. The published incident record is thin, for the disclosure reasons set out in the Introduction, and synthetic red-team exercises test what the exercise designer imagined. Organisations deploying transition-state detection are therefore in the position of tuning a control whose true positive rate they cannot measure, and the honest response is to instrument heavily, retain the traces, and expect the first genuine detection to reveal that the model was wrong in ways nobody anticipated.

A related measurement question concerns what counts as a detection at all. In conventional intrusion detection the event is discrete: a payload was delivered, a process spawned, a connection made. Agentic compromise is gradual in several of its forms, particularly the memory injection pattern, where the moment of compromise and the moment of damage are separated by weeks and neither is individually anomalous. A detection programme measured on time-to-detect will report poorly against attacks

whose beginning is invisible and whose end is unremarkable, and the metric will drive investment toward the fast, loud cases rather than the slow, expensive ones.

There is also an interaction between this control and the platform team's development practice that is worth anticipating. Every change to an agent's tools, prompts, or retrieval configuration alters its behavioural profile, and a detection system that has not been told about the change will interpret the new behaviour as anomalous. Teams that deploy agent changes several times a week will therefore generate a burst of false positives after every release unless the release pipeline signals the detection layer. Wiring that signal is a small piece of integration work, it belongs in the deployment pipeline rather than in the security platform, and its absence is the most common reason transition-state detection gets switched off within a quarter of being deployed.

A note on how this interacts with the enforcement controls of Chapter 8, since the two are frequently confused in design discussions. Transition modelling learns what an agent does and flags departures from it. Scope declaration states what an agent is permitted to do and refuses anything outside it. The first is descriptive and the second prescriptive, and they fail in opposite directions: the descriptive model permits anything the agent has done before, including things it should never have been doing, while the prescriptive declaration refuses anything nobody thought to write down, including legitimate work. Running both means an action must be within declared scope and consistent with observed behaviour, which is a narrower gate than either alone and is the arrangement mature deployments converge on.

Where the two disagree, the disagreement is itself informative and should be routed somewhere a human will read it. An action inside declared scope that the agent has never performed suggests the declaration is broader than the agent's real requirements, which is a permission to narrow. An action outside declared scope that the agent performs routinely suggests the declaration is wrong, which is a specification to fix. Neither case is an incident, and both are the kind of finding that improves the system's

security posture between incidents, which is where most durable improvement actually comes from.

A closing note on where this control should sit. Transition-state analysis can run as detection, producing alerts for human review, or as enforcement, blocking the anomalous transition before it executes. Detection is cheaper, safer to deploy, and useless against an attack that completes in four hundred milliseconds. Enforcement stops the action and will occasionally stop legitimate work, with consequences that depend entirely on what the agent was doing. The workable arrangement runs enforcement on the small set of transitions leading to irreversible actions and detection on everything else, which bounds both the operational risk and the false-positive cost. The next chapter examines the same detection problem on the internal communication surface, where the volume is higher and the visibility, according to the survey data, is close to absent.

Chapter 8: Monitoring Agent-to-Agent Communication

Chapter 6 described the security properties of inter-agent communication and argued that the trust assumptions inherited from service architecture do not hold. This chapter takes up the operational consequence, which is that most organisations cannot see this traffic at all, and describes the control that has emerged to address it.

8.1 The Agent Communication Visibility Gap

Survey work across enterprise security functions found that only 24.4 percent of organisations maintain active visibility into the communication pathways between agents in their environments. The remaining three quarters operate with what has come to be called dark agent traffic: workflows that exchange context, delegate authority, and pass credentials to one another entirely outside the collection scope of the security operations function.

Visibility into agent-to-agent communication pathways

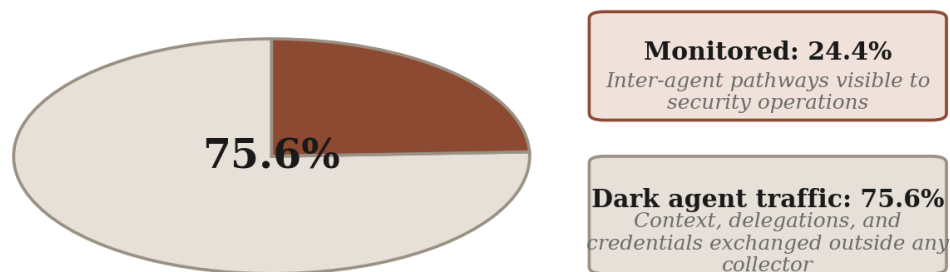


Figure 8.1 Reported visibility into inter-agent pathways. Three quarters of the population cannot observe the traffic described in Chapter 6.

That figure is worse than the comparable number for any other traffic class an enterprise carries, and understanding why requires looking at how this traffic came to exist. Network traffic is visible because the network is instrumented and has been for thirty years. Application traffic is visible because application performance monitoring became standard practice. Inter-service traffic became visible when service meshes

arrived and made it a property of the platform. Inter-agent traffic has no equivalent history. It travels over connections the orchestrator establishes, in formats the orchestration framework defines, between components that the network layer sees as one application, and nobody built a collector for it because until recently the traffic did not exist.

Several properties of this traffic make retrofitting visibility harder than it was for earlier classes.

Much of this traffic is internal to a process boundary in many deployments. Where an orchestrator manages several agents within a single runtime, messages between them are function calls rather than network events, and no amount of network instrumentation will reveal them. Collection has to happen inside the orchestrator, which means the orchestrator has to be built or configured to emit it.

Topology is generated at runtime. A planner that decomposes an objective decides which sub-agents to create and how they will communicate, and that decision reflects the objective rather than a design document. A network policy or monitoring configuration written against the expected topology describes a system that may not be the one running. This is the inventory problem from Chapter 1 reappearing on the communication surface, and it has the same answer: the model has to be derived from runtime telemetry rather than maintained by hand.

Volume is high and the content is dense. A single user request may generate hundreds of inter-agent messages, each carrying natural language rather than structured fields. Storing them all is expensive, and evaluating them all is more expensive still, which forces a sampling or tiering decision that has security consequences. Sampling inter-agent traffic means accepting that some fraction of the delegation path will be invisible after the fact, and that fraction is exactly where an investigator will need to look.

Encryption between agents complicates collection in a way that surprises teams who did the transport security work first. A well-engineered deployment establishes mutual TLS

between components, which is correct and which also means that a passive network collector observes ciphertext. Observing the content requires either terminating the connection at a mediating component, which is the semantic firewall pattern described in section 8.2, or instrumenting the endpoints to emit what they sent. Both are legitimate, both cost something, and neither is what an organisation gets by pointing a network sensor at the cluster. The relevant point for planning is that transport security and message visibility are separate pieces of work, and completing the first does not advance the second.

Ownership of this traffic is genuinely ambiguous in a way that slows remediation, and naming the ambiguity helps. Network teams regard it as application traffic and outside their remit. Application teams regard it as platform behaviour they did not write. Platform teams regard the messages as application content they should not be inspecting. Security teams regard it as somebody else's telemetry to emit. Each position is defensible, the result is that nobody instruments it, and the resolution is an explicit assignment rather than a technical decision. In organisations that have closed this gap, the assignment usually went to whichever team owns the orchestration platform, on the reasoning that they are the only party positioned to emit the data at all.

Vendor-managed agentic services introduce a harder version of the same problem and deserve separate consideration during procurement. Where an organisation consumes agentic capability as a service, the inter-agent communication happens inside the provider's infrastructure and the customer receives whatever telemetry the provider chooses to expose. That is frequently a summary of tasks completed, with no delegation records, no message content, and no reasoning traces. The security question to ask before signing concerns what the customer can observe, rather than how the provider describes its own controls, because an organisation that cannot observe cannot investigate and will be answering for an incident it has no data about.

Consequences of the gap become clearest during investigation. Chapter 6 described cascading injection, in which a payload entering at an edge agent propagates inward to a

component with consequential authority. Investigating such an incident requires the message chain: what the edge agent received, what it forwarded, what each subsequent hop passed on. An organisation without inter-agent visibility possesses the entry point, meaning the original customer message, and the outcome, meaning the unauthorised write, with nothing in between. The investigation stops at the observation that something occurred, which is where the 247-day detection figure from Chapter 1 comes from.

Recursive agent creation makes the gap wider than the survey figure implies. When an orchestrator spins up sub-agents, and those sub-agents spin up their own, the chain of command becomes several layers deep and each layer was created by software rather than by a person. An attacker who compromises a component partway down that chain operates in a region of the system that no human has ever inspected. Whether the survey respondents who reported having visibility have visibility into these dynamically created layers, or only into the top-level agents they configured, is a question the survey does not resolve, and practitioner accounts suggest the latter is common.

The regulatory position adds urgency that is separate from the security argument, and organisations in supervised sectors should register it now rather than when an examiner raises it. Requirements to demonstrate control over data flows, to evidence that access was appropriate, and to reconstruct decision paths all assume that the flows in question are observable. An institution whose agents exchange customer context over unmonitored internal pathways cannot evidence what was shared with what, which is a gap that will be found. Chapter 12 develops this and the practical point belongs here: the visibility work has a compliance justification available to any team that finds the security case difficult to fund.

Cost is the objection that surfaces first when this work is proposed, and it is worth answering directly. Full capture of inter-agent traffic in a busy deployment produces a data volume comparable to the deployment's entire application logging, and organisations paying per gigabyte will reject the proposal on sight. The answer is that

full capture is not required for most of the benefit. Recording the delegation graph, meaning which agent asked which other agent to do what, at what time, under which session, costs a small fraction of recording the message bodies and answers most investigative questions about propagation. Message bodies can be retained selectively, for delegations that cross a trust boundary or target a consequential capability, which is the same tiering the enforcement control uses.

Before describing the control, one observation about scope is worth making. The visibility problem here is a specific instance of a general one: security monitoring has always been placed where the architecture put its boundaries, and agentic architectures put their most interesting boundaries in places that no monitoring convention covers. The boundary between an agent and the untrusted content it retrieves, the boundary between one agent and another, and the boundary between the reasoning that produced an action and the action itself are all security boundaries with no established collection point. Chapter 9 addresses the third of these. This chapter addresses the second.

8.2 Implementing Semantic Firewalls

Mediating components positioned in the inter-agent communication path have emerged as the control for this surface, generally described as semantic firewalls. The term invites a comparison with network firewalls that clarifies the difference in what is being evaluated. A network firewall decides on the basis of addresses, ports, and protocol conformance, all of which are structural properties available by inspection. A semantic firewall decides on the basis of what a message asks for, which requires interpretation.

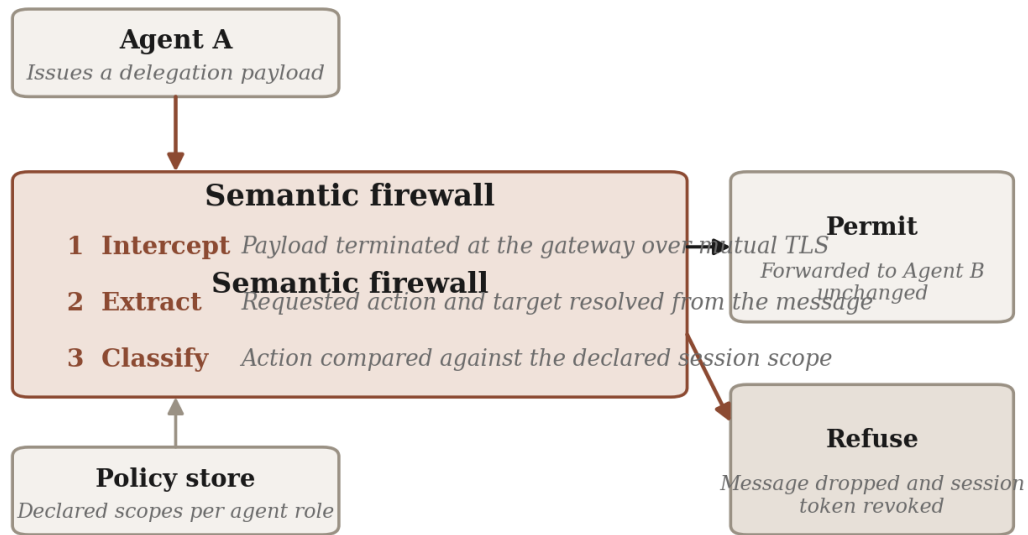


Figure 8.2 A semantic firewall in the delegation path, evaluating requested action against declared scope before forwarding.

Operationally the component performs four steps. It terminates the connection, so that the message is available rather than merely observed, which in a mutually authenticated environment means the firewall holds credentials permitting it to act as an endpoint for both parties. It parses the payload and extracts the requested action, the target resource, and any authority the message claims. It evaluates that extracted intent against the policy applicable to the session, meaning the scope declared when the originating task began. And it permits, refuses, or escalates the message according to the result, with a refusal accompanied by revocation of the sending agent's session credential where the policy calls for it.

Intent extraction is the part that carries the engineering risk, and it should be understood as approximate. A message can request something in many ways, and an evaluator that examines phrasing can be misled by phrasing. This is the same limitation identified for input filtering in Chapter 2, and it has the same implication: a semantic firewall that decides purely by interpreting message content is a probabilistic control placed in an enforcement position.

A design that avoids this weakness reframes the evaluation, and the reframing is the most important idea in this chapter. Rather than asking whether the message is malicious, which requires judgement about content, the firewall asks whether the action the message requests falls within the scope declared for the session, which requires a lookup. The first question has no reliable answer. The second has a definite one, provided the session declared its scope at the outset and the firewall can retrieve it. A message requesting a bulk export within a session scoped to a single customer record is refused, and the refusal does not depend on recognising that the message was hostile.

Content evaluation retains a role in this design, and it is a subordinate one. Where the requested action is within scope, a classifier examining the message for injection patterns provides a second opinion whose errors are tolerable, because a false negative means an in-scope action proceeds rather than an out-of-scope one. Where the action is outside scope, the classifier is not consulted, because the answer is already determined. This ordering, deterministic check first and probabilistic check second, is what keeps a model-based component out of the position where its failures are consequential.

Tiering follows from the cost analysis in Chapter 6. Full evaluation of every inter-agent message multiplies inference cost by the density of the communication graph, which is unaffordable in a system doing real multi-agent work. The tiers that have worked in deployment apply a scope check to every message, since it is a lookup and costs microseconds; apply classifier evaluation to messages crossing from a component exposed to external content into one that is not; and apply full semantic evaluation only to messages that delegate authority or that target an agent holding irreversible capability. Placing the boundaries requires knowing which agents hold which authority, which is the identity work of Chapter 3, and it is the fourth control in this book to depend on that foundation.

Policy authorship is the question that determines whether this control is maintainable, and it is where deployments most often stall after the technology is working. Somebody has to state, for each agent role, which actions fall inside its scope. Security teams

generally lack the domain knowledge to write these declarations, since they do not know what a claims processing agent legitimately needs. Business teams have the knowledge and no incentive to constrain themselves. Platform teams have neither the knowledge nor the mandate. The arrangement that works assigns authorship to the team that owns the agent, review to security, and a default-deny posture on anything undeclared, which converts the incentive problem into a straightforward one: the team that wants its agent to work has to describe what it does.

A useful accelerant is to derive an initial declaration from observed behaviour rather than asking anyone to write one from scratch. Running the agent in observation mode for a period produces a record of which tools it called against which resources, and that record is a first draft of its scope. The team then reviews the draft, removes anything that should not have happened, and adopts the remainder. This inverts the usual painful sequence, in which a team is asked to specify behaviour it has never enumerated, and it has the additional benefit of surfacing the unexpected access that section 7.2 describes the transition model finding.

Three deployment considerations deserve mention because they determine whether the control survives contact with a platform team.

Latency sits on the path of every inter-agent message, and a multi-agent task involving fifty messages inherits fifty times the per-message cost. A scope check adds microseconds and is invisible. A classifier adds milliseconds and is noticeable. A full model evaluation adds hundreds of milliseconds and will be removed by the first team whose task latency doubles. The tiering above is a latency budget as much as a cost budget.

Availability makes the firewall a dependency of the entire agentic estate, and its failure mode has to be decided in advance rather than during an incident. Failing open preserves availability and removes the control precisely when an attacker who can disrupt the component would want it removed. Failing closed preserves the control and converts a firewall outage into an application outage. The defensible position is to fail

closed for messages targeting irreversible capability and open for the rest, which requires the same classification the tiering already needs.

Data sensitivity is the consideration teams most often overlook. A semantic firewall terminates every inter-agent message, which means it sees every piece of context those messages carry, including customer data, credentials in transit, and the full reasoning of the systems it mediates. It is among the most sensitive components in the environment, and its own access controls, retention policy, and log handling deserve the scrutiny given to a production database. Deploying a security control that becomes the most attractive target in the estate is a familiar pattern and remains an avoidable one.

Two failure modes have been observed in early deployments of this control, and both are worth designing against explicitly. The first is scope inflation. When a firewall refuses an in-scope-looking action, the fastest resolution available to the affected team is to widen the declared scope, and repeated widening produces declarations broad enough that nothing is ever refused. The control remains in place, generates no alerts, and provides no protection, which is worse than its absence because it appears on the control inventory. Guarding against this requires reviewing scope declarations as a governed artefact, with changes recorded and periodically examined for drift, rather than letting them be edited by whoever is on call.

The second is evaluation of the wrong artefact. A firewall placed at the network boundary between orchestrator instances sees the messages that cross that boundary and misses the ones exchanged inside a single instance, which in many deployments is the majority. Teams have deployed this control, confirmed it was passing traffic, and only later discovered that the delegations they most wanted to observe never crossed the wire. Establishing where the messages actually travel, before selecting a placement, is a prerequisite that sounds obvious and is regularly skipped.

A question that arises in review is whether this control could be implemented inside the orchestrator rather than beside it, which would remove the latency of an additional hop and the availability risk of an additional component. It could, and several orchestration

frameworks are moving in that direction. The objection is one of separation: a control implemented inside the component it constrains is disabled by whoever compromises that component, and the orchestrator is a plausible target precisely because it holds the planning loop. The defensible arrangement places policy evaluation in a component with its own identity and its own failure domain, even where that costs a hop, and treats orchestrator-native controls as a useful additional layer rather than as a replacement.

Measuring whether this control is working requires care, because its success metric is a non-event and the natural proxies are misleading. A firewall that refuses nothing may indicate an environment with no attempted violations or a declaration so broad that nothing can violate it, and the two are indistinguishable from the refusal count alone. Useful instrumentation therefore records near-misses as well as refusals: actions that fell inside scope but only just, actions that required a widening of scope to proceed, and sessions whose declared scope was unusually broad at the outset. The distribution of those measurements tells an operator whether the control is doing work, which the refusal count alone does not.

Periodic adversarial testing supplies the other half of the assurance, and it is cheap to arrange for this control specifically. Constructing a test is a matter of instructing an agent, through a legitimate channel, to attempt something outside its declared scope and confirming that the attempt is refused and recorded. This requires no exploit development and no specialist tooling, it can be automated and run continuously, and it verifies the property that matters. Organisations that run this as a scheduled check catch the configuration drift that silently disables the control, which is a more common failure than any adversary defeating it.

A final observation about the relationship between this control and the previous chapter. Transition-state detection watches what an agent does; a semantic firewall watches what agents ask of each other. The two are complementary, and they share a prerequisite, which is a declared statement of what each session is supposed to accomplish. Without that declaration, transition modelling has no notion of intended

behaviour to compare against and the firewall has no scope to evaluate against. Both controls therefore push the same requirement back onto the platform: the orchestrator must state, in machine-readable form and at session start, what the agent has been asked to do. Chapter 9 turns to the telemetry that makes all of this observable in the first place.

Chapter 9: Building the Agentic SOC

The controls described in the two preceding chapters are only as good as the data they operate on, and the data required does not exist in a conventional security operations environment. This chapter covers what has to be collected, how it should be structured, and how incident response changes when the compromised asset is a process that will not exist by the time an analyst opens the ticket.

9.1 Ingesting Agent Telemetry

A security information and event management platform is a correlation engine over a corpus of events, and its usefulness depends entirely on whether that corpus contains the events that matter. For agentic workloads the corpus in most organisations does not. Infrastructure logs record that a container started. Network logs record that a connection was made. Application logs record that an API returned a status code. None of them records what the agent was trying to do, why it decided to do it, or what it read before deciding.

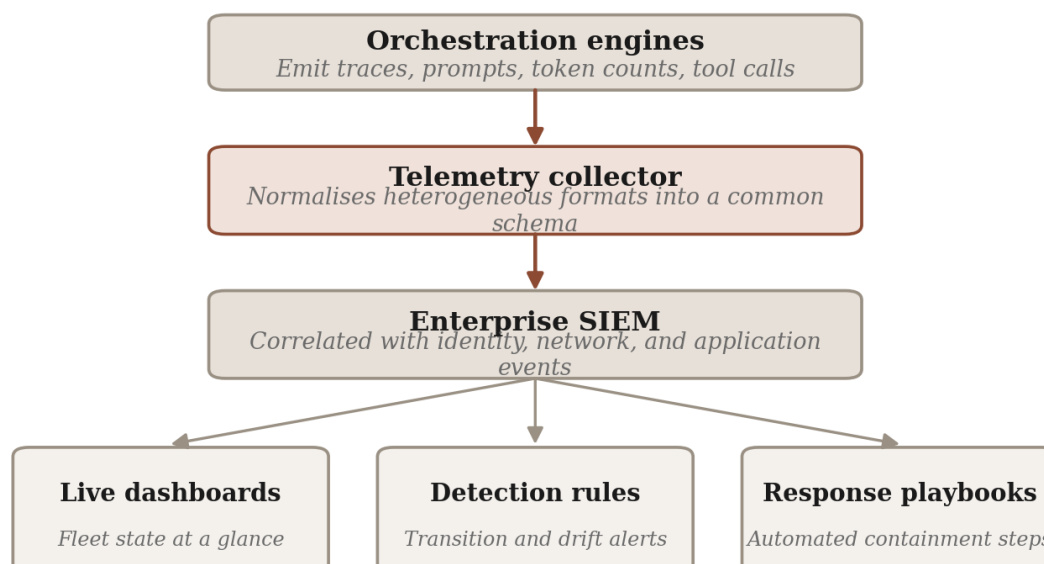


Figure 9.1 The agentic telemetry path. Orchestrators are the only components that hold the reasoning context, which makes them the collection point.

Four categories of telemetry have to be collected from the orchestration layer, because the orchestrator is the only component in the stack that holds them.

Prompt traces are the complete lineage of a task: the objective as received, the plan the model produced, the intermediate reasoning at each step, the observations returned by tools, and the revisions the plan underwent. This is the record that explains why an action was taken, and it has no analogue in conventional systems because conventional systems have no reasoning to record. Prompt traces are also the largest of the four by volume and the most sensitive by content, since it contains everything the agent read.

System prompt state records the instructions governing the agent's behaviour and, more importantly, every change to them. A modification to a system prompt is a change to what the software does, comparable in significance to a code deployment, and in most organisations it is made through a console by whoever has access, without review, versioning, or an audit record. Capturing prompt state and treating its modification as a change event brings this under the same discipline as any other production change.

Token metrics record input, output, and cached consumption per step, per session, and per agent. Their detection value was covered in Chapter 7 and their financial value in Chapter 5, and they are the cheapest of the four to collect because the orchestrator already counts them for billing.

Tool-call logs record which tool was invoked, with which parameters, against which target, with what result. Many orchestration frameworks emit these by default, which makes them the usual starting point for a telemetry programme. Their limitation is that they record the action without the reasoning, so an investigator can establish what happened and not why, which is sufficient for reconstruction and insufficient for determining whether an injection occurred.

Telemetry category	What it answers	Collection difficulty
Prompt traces	Why the agent chose this action	High volume, high sensitivity, often disabled by default

Telemetry category	What it answers	Collection difficulty
System prompt state	What the agent was instructed to be, and when that changed	Low volume, requires treating prompt edits as change events
Token metrics	How much work was performed and whether it was progressing	Already counted for billing; usually the easiest to obtain
Tool-call logs	What the agent actually did to which system	Commonly emitted by default; lacks reasoning context
Delegation records	Which agent asked which other agent for what	Frequently absent; requires orchestrator instrumentation

Table 9.1 Agentic telemetry categories, the investigative questions each answers, and the practical obstacles to collecting it.

A fifth category belongs on this list and is absent from most collection programmes because no framework emits it by default. Delegation records capture which agent requested work from which other agent, under which session, with what claimed authority, and they are the telemetry that makes the propagation chains of Chapter 6 reconstructible. Where an orchestrator manages sub-agents within a single process, these interactions leave no network trace and appear in no log unless the orchestrator is configured to record them. Organisations running multi-agent workloads without delegation records can observe the entry and exit of a cascade and nothing of its path.

Deciding what to collect is easier when the question is framed by the investigation it has to support rather than by what the platform happens to emit. The reconstruction an analyst will need to perform runs: which principal authorised this session, what objective was declared, what did the agent read, what did it decide, which tools did it call, what did those tools return, which other agents did it involve, and what changed in any system as a result. Each clause in that sentence maps onto a telemetry category, and a collection programme can be assessed simply by asking which clauses it can currently answer. Most programmes answer the sixth and the last, and nothing else.

Three problems attend this collection, and each has a workable answer that costs something.

Volume is the first. Prompt traces for a busy deployment can exceed the total log volume of the rest of the environment, and organisations paying for ingestion by the gigabyte will notice immediately. The approach that works is tiered retention: full traces held for a short window measured in days, structured summaries retained for months, and complete traces preserved indefinitely for sessions flagged by any detection rule. The window has to exceed the time an incident typically takes to surface, and given the detection figures in Chapter 1, a window measured in days is an acknowledged compromise rather than a solution.

Sensitivity is the second. A prompt trace contains everything the agent read, which in a customer-facing deployment means customer data, and in an internal deployment may mean source code, financial records, or personnel information. The telemetry store therefore inherits the classification of the most sensitive data any agent handles, and it aggregates across agents, which makes it more sensitive than any individual source. Redaction at collection helps and reduces investigative value, since the redacted field is sometimes the one that mattered. Most organisations that have worked through this apply strict access controls and full audit on the trace store rather than redacting, treating it as a crown-jewel system.

Format heterogeneity is the third. Every orchestration framework structures its traces differently, and an organisation running several has as many schemas. Normalisation into a common representation is unglamorous engineering that has to be redone whenever a framework updates, and it is the work that most often gets deferred. Deferring it means each framework's telemetry is queried separately, which defeats correlation, which is the entire purpose of the platform.

Retention policy deserves a paragraph of its own because the default is wrong in a specific and expensive way. Observability platforms are configured for operational debugging, where the useful window is hours and anything older is noise, and agentic telemetry ingested through an observability pipeline inherits that default. The security window is set by the detection interval, which the figures in Chapter 1 put at roughly

eight months for this class of compromise. A deployment retaining traces for seven days has telemetry that will answer questions about performance and none of the questions an investigator will ask. Reconciling these two windows without paying for eight months of full traces is what the tiered approach above is for, and the reconciliation has to be a deliberate decision rather than an inherited default.

Sampling requires similar care. Observability sampling discards a proportion of traces at random, which preserves statistical properties and is entirely appropriate for measuring latency distributions. Applied to security telemetry it discards evidence, and the discarded trace is as likely as any other to be the one that mattered. Where volume forces sampling, the selection has to be biased toward security relevance: retain every session that touched a consequential tool, every session flagged by any detection rule, and every session whose scope was widened mid-execution, then sample the remainder. This is a different sampling policy from the one the observability team will configure, and the difference has to be explained rather than assumed.

Instrumentation of the model provider boundary is the collection point teams most often overlook, and it carries information available nowhere else. Every call to a provider carries the assembled context the agent constructed, which is the closest thing to a complete record of what the model was asked to reason about. Capturing at this boundary has the practical advantage of being framework-independent, since it observes the outbound API call rather than the orchestrator's internals, and it therefore survives a change of orchestration framework. Where the cost mediation proxy of Chapter 5 has already been deployed, it sits exactly at this boundary and can emit the telemetry as a side effect of work it is doing anyway.

A standardisation caution belongs here as well. Several efforts are under way to define common schemas for agentic telemetry, and adopting one early carries the usual risk of betting on a specification that does not prevail. The mitigating position is to collect a superset of what any candidate schema requires, store it in a form that preserves the original structure, and treat the schema as a projection applied at query time rather than

at ingestion. This costs storage and preserves the option to re-project when the standard settles, which is a better trade than discarding fields that turn out to matter.

Correlation is where the value is realised, and the correlations that matter cross domains. A prompt trace joined to an identity event answers which principal authorised the session. A tool-call log joined to a database audit record confirms whether the action the agent believed it performed is the action the system recorded. A token metric joined to a billing record attributes cost to a task. A delegation record joined to a network event reconstructs a cascade. None of these is available to an organisation that collects agentic telemetry into a separate system, which is the arrangement that observability vendors currently encourage and that security teams should resist.

9.2 Designing SOC Playbooks

Incident response procedures assume an asset. Isolate the host, image the disk, preserve memory, block the address, disable the account. Every step in a standard playbook operates on something that persists long enough to be operated on. An agentic incident presents none of these. The compromised component is an execution context that may have terminated before the alert was delivered, running in a container that no longer exists, having acted through credentials that expired on their own.

Agentic incident response: three phases, minutes not hours

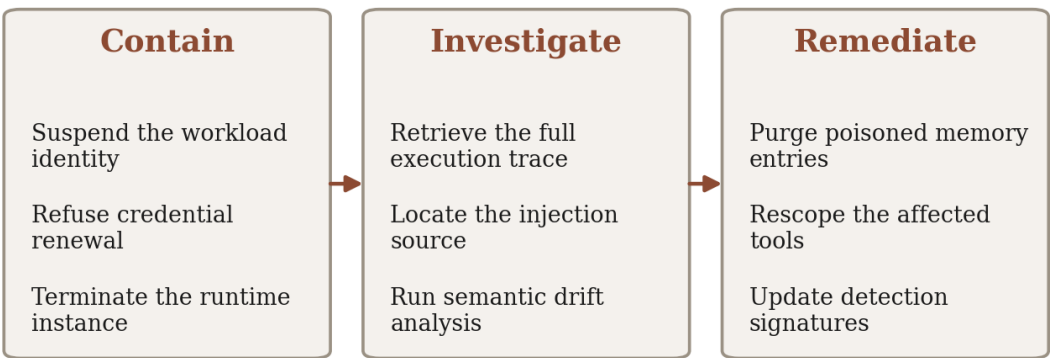


Figure 9.2 Response phases for an agentic incident. Containment operates on identity and runtime rather than on a device.

Containment operates on two levers, and both must be reachable by an analyst without engineering assistance. The first is the workload identity: suspending it at the issuance layer causes the next credential renewal to fail, which stops the agent within one credential lifetime. This is why the short lifetimes advocated in Chapter 3 matter operationally as well as defensively; an organisation using long-lived static credentials has no equivalent lever and must revoke a key that other agents also use. The second is the runtime: terminating the execution environment stops the current activity immediately, which matters when the credential lifetime is longer than the damage window.

Both levers should be exercised in a test before they are needed. Teams that have run this exercise frequently discover that suspending a workload identity requires a platform engineer with console access, that no runbook documents which console, and that the agent in question was deployed by a business unit whose on-call rotation does not exist. Finding that out at two in the morning during an active incident is the avoidable failure that this paragraph exists to prevent.

Investigation begins with the execution trace and proceeds backward. The immediate question is where the deviation entered, and semantic drift analysis provides the answer efficiently: the step at which consecutive similarity broke discontinuously is where the agent's objective changed. Reading the context available at that step usually identifies the injected content and its source, which then determines the scope of the incident. If the source is a document in a shared repository, every agent that read it is in scope. If it is a customer ticket, every agent processing tickets from that channel is in scope.

Scoping must extend backward past the observed event, for the reason set out in Chapter 2. If the incident involved retrieved memory, the relevant window begins when the poisoned entry was written, which may precede the incident by months. Provenance on memory entries makes this a query. Without provenance, the investigator faces a

choice between assuming the memory is clean and rebuilding it entirely, and neither option is satisfactory.

Remediation has three components that map onto the three places an agentic compromise can persist. Poisoned memory entries must be purged, which requires knowing which entries derived from the compromised session and is straightforward where provenance was recorded. Tool scopes must be reviewed and narrowed where the incident showed an agent holding authority it did not need, which is the point at which an organisation converts an incident into a permanent reduction in exposure. And detection content must be updated so that the specific pattern is caught next time, which for transition-state models means adding the observed transition to the prohibited set rather than writing a signature.

Preservation requires attention before containment, which inverts the usual order and is easy to get wrong under pressure. Terminating a runtime destroys whatever state existed only in memory, which for an agent includes the working context that explains its current behaviour. The correct sequence captures the execution state first, then terminates, and the capture has to be a single command an analyst can run rather than a procedure requiring engineering support. Where the platform does not support state capture, the compensating control is to make certain the trace pipeline is emitting continuously, so that terminating the runtime loses only the final seconds.

Communication during an agentic incident has a feature that no other incident class presents, and response leads should be prepared for it. When asked what happened, the agent will answer, fluently and with apparent confidence, citing whatever context it holds. Chapter 2 described why that answer is unreliable when memory has been poisoned. The operational hazard is that the answer is also persuasive, and it will circulate: an engineer asks the agent, receives a coherent explanation, and reports it to the incident channel, where it becomes the working theory before anyone has read a trace. The procedural remedy is to state in the playbook that agent self-reports are not

evidence, which sounds obvious in writing and is regularly forgotten at three in the morning.

Two categories of response that work well elsewhere fail here and are worth naming so that playbooks do not include them. Blocking an address accomplishes nothing, because the attacker interacted with a data source rather than with the target's network and has not been connected since. Disabling a user account accomplishes little where the agent holds its own identity, and disables a person's access where it does not, which produces a business disruption without stopping the agent.

Three organisational changes support all of the above, and they are harder than the technical work.

The response window is compressed. An agent acting at machine speed can complete a damaging sequence in the time it takes an analyst to open a ticket, which means containment for high-consequence agents has to be automated rather than triaged. Automating containment means accepting that some legitimate work will be stopped by mistake, and the organisation has to decide in advance which agents warrant that trade rather than deciding during an incident.

Analyst skills have to broaden. Reading an execution trace and judging whether an agent's reasoning was manipulated is a different competence from reading a packet capture or a process tree. It resembles the work of a fraud analyst assessing whether a sequence of transactions tells a coherent story, and organisations have had more success developing this capability internally than hiring for it, since the external market for it barely exists.

Ownership has to be settled before an incident rather than during one. The agent belongs to a business unit, the orchestrator to a platform team, the model to an external provider, the data to whoever owns the system it came from, and the incident to security. Every one of these parties has to act during the response, and none of them reports to the others. Organisations that have handled agentic incidents well established

this map in advance, and those that have not describe the first hour of the incident as an exercise in finding out who is permitted to stop the thing.

Exercising these procedures deserves the same treatment as any other response capability, and the exercise design differs from what security teams are used to running. A tabletop for an agentic incident is more useful than usual, because the questions it surfaces are organisational rather than technical: who can stop this, which system holds the trace, who owns the memory store, what happens to the business process while the agent is suspended. Most organisations that run this exercise for the first time fail it, in the sense that the group cannot answer several of those questions, and the failure is the point. The technical simulation matters less here than the discovery of which authority is missing.

A specific exercise worth running is the reconstruction drill. Take a completed agentic task from the previous week, chosen at random, and ask the team to reconstruct from telemetry alone what the agent was asked to do, what it read, which tools it called, what it wrote, and which principal authorised it. The exercise takes an afternoon, requires no scenario design, and produces a precise inventory of what the telemetry programme is missing. Teams that can complete it have the data an investigation requires. Teams that cannot have learned exactly which gap to close first, which is more actionable than any maturity assessment.

A final point about the relationship between this chapter and the rest of the book. Every control described in Parts I through III has assumed the ability to observe what agents do, and the observation is not free, not automatic, and not a property of any platform an organisation is likely to have purchased. Telemetry is the foundation on which identity attribution becomes verifiable, behavioural modelling becomes possible, and incident response becomes something other than guesswork. An organisation with limited capacity to invest should build the telemetry first, because every subsequent control depends on it and none of them can be retrofitted onto data that was never collected.

Automation of the response itself deserves a closing note, since the compressed response window argued for above pushes in that direction and the risks are real. An automated containment action that suspends an agent identity on a detection is fast enough to matter and will occasionally suspend a business-critical workflow on a false positive. The mitigations are ordinary and effective: apply automation only to agents whose interruption is tolerable, require a second corroborating signal before automated action on the rest, and make reinstatement a one-step operation so that a mistaken suspension costs minutes rather than an escalation. What should be avoided is the arrangement where automation is enabled globally because it was easier to configure that way, which is how a detection tuning error becomes an outage.

There is a further consideration about the pace at which any of this can be introduced. The controls in this part are not independent, and attempting them in the wrong order wastes effort. Telemetry precedes everything. Identity attribution, from Chapter 3, is required before telemetry can be attributed to anything. Behavioural modelling requires attributed telemetry with enough history to learn from. Enforcement requires a model to enforce against and a declaration to compare with. An organisation that begins with enforcement, which is the layer with the clearest vendor offering, will find it has bought a component with nothing to feed it.

This part has argued that detection for agentic systems has to be rebuilt around behaviour and equipped with telemetry that conventional environments do not produce. The controls described are technical, and each of them has turned out to depend on a declaration that the platform team must supply and on an identity foundation that the security team must build. Part IV turns to the governance structures that make those obligations stick, beginning with the application of zero-trust principles to the execution of tools by autonomous software.

Part Four: Cloud Governance and Trust

Chapter 10: Zero-Trust Governance Frameworks

The preceding parts described what agentic systems expose and how an organisation might observe them. This part addresses the structures that make observation and constraint durable: the architectural principles, the emerging standards, and the regulatory obligations that will determine what a defensible deployment looks like. The three chapters progress from what an organisation can build for itself, through what the standards bodies are converging on, to what regulators and cloud contracts will require of it.

Zero trust is the natural starting point because its central premise happens to describe the agentic problem with unusual precision. The principle, articulated in NIST Special Publication 800-207 and elsewhere, holds that no implicit trust should be granted on the basis of network position, asset ownership, or prior authentication, and that every access decision should be made on the strength of current evidence. Enterprises have spent a decade applying this to human users and to services. Applying it to software that writes its own instructions requires extending the model in two directions, which this chapter takes in turn: isolating the environment in which agent-generated code runs, and moving authorisation from the beginning of a session to every step within it.

10.1 Sandboxed Tool Execution

An agent that can execute code is the most useful configuration available and the most dangerous. Code execution allows an agent to transform data, run analyses, drive systems that expose no clean interface, and correct its own errors, which covers a large share of what makes agentic automation valuable. It also means the system is running instructions that no person reviewed, composed moments earlier by a component that can be influenced by any text it read.

The zero-trust position on this is straightforward once stated. Any environment in which an agent executes generated code must be treated as compromised by default, because there is no mechanism by which the organisation can establish that it is not. Static

review is impossible, since the code did not exist at review time. Behavioural analysis of the code before execution is possible in principle and defeated by the same paraphrase problem that defeats content filtering. The remaining option is containment: permit the execution and constrain what its consequences can be.

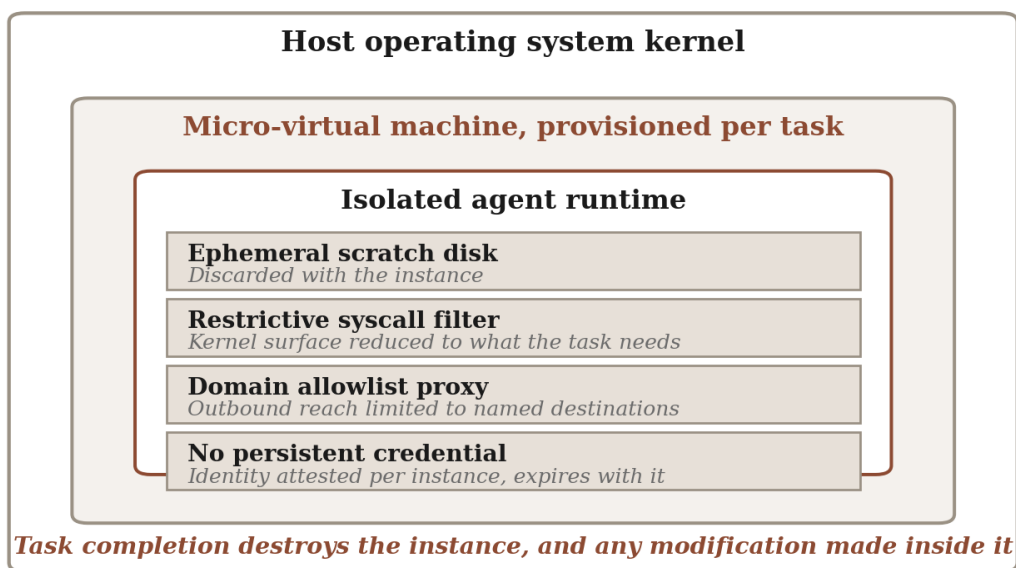


Figure 10.1 Layered isolation for agent code execution. Each layer constrains a different escape path, and the instance is destroyed on completion.

Production architectures implement this with lightweight virtualisation. Micro-virtual machines, of which Firecracker is the most widely deployed example, provide hardware-assisted isolation with startup times measured in tens of milliseconds, which makes per-task instantiation practical in a way that conventional virtual machines never were. The distinction from container isolation matters here. Containers share a kernel, so a kernel vulnerability is an escape path, and container escapes are a recurring feature of the vulnerability record. A micro-virtual machine presents a separate kernel and a much smaller interface to the host, which raises the cost of escape by a wide margin.

Four constraints operate inside the instance, and each addresses a distinct failure.

Ephemeral storage means the instance receives scratch space that is discarded when it terminates. Anything the agent writes, including anything an attacker persuaded it to

write, disappears with the instance. This removes the persistence mechanism that conventional intrusion depends on, and it also removes a legitimate capability, since an agent that needs to preserve output must write it to a designated external location where the write is observable.

System call filtering restricts the kernel interface available to the runtime, using `seccomp` or an equivalent mechanism, so that the set of operations a compromised process can attempt is narrowed to those the task requires. The effect is to shrink the attack surface against the host to a fraction of what a general-purpose runtime exposes.

Egress control constrains what the instance can reach on the network. This is the constraint that matters most in practice and is most often configured loosely. An agent runtime with unrestricted outbound access can exfiltrate anything it holds to any destination, and a proxy enforcing an allowlist of named destinations turns that from a certainty into a bounded set. The operational objection is that agents legitimately need to reach varied endpoints, which is true and is an argument for maintaining the allowlist rather than for abandoning it.

Identity scoping applies the delegation model of Chapter 3 to the runtime instance itself. The instance attests its own state and receives a credential valid for the task and for the instance lifetime, which means a credential extracted from a sandbox expires with the sandbox and authorises only what the task required. An instance holding a long-lived credential has defeated most of the value of the isolation around it, since the credential outlives the containment.

Termination is the property that ties these together and deserves stating explicitly because it is easy to compromise for convenience. When the task completes, the instance is destroyed, and every modification made within it goes with it. Deployments that reuse instances across tasks, for the entirely reasonable purpose of avoiding startup cost, have reintroduced persistence and with it the possibility that one task's compromise affects the next. Where reuse is necessary for performance, the pool should be recycled on a

short cycle and instances should never be shared across sessions belonging to different principals.

Three implementation realities complicate this picture and are worth stating plainly, since the design carries costs that its clean description conceals.

The first is that isolation constrains capability in ways users notice. An agent that cannot reach an unlisted endpoint cannot complete a task requiring one, and the resulting failures arrive as support tickets attributing the problem to security. Maintaining the allowlist is continuing operational work with an owner, and deployments where nobody owns it decay toward either a permissive configuration or a stream of unexplained failures.

The second is that instance startup, even at tens of milliseconds, is a cost on every task, and for workloads composed of many short tasks the aggregate is significant. Teams respond by pooling, which as noted weakens the isolation, and the correct response is to size the pool policy against the sensitivity of the work rather than against the latency budget alone.

The third is that isolation addresses the consequences of compromised execution and does nothing about compromised decisions. An agent manipulated into performing a destructive but entirely authorised action, using a legitimate tool through its normal interface, is not constrained by any of this. The sandbox contains code execution; it does not contain a payment instruction. This is why the authorisation work in the next section is a separate control rather than a refinement of the same one.

A further constraint belongs in this list and is absent from most deployments because it addresses a case that feels theoretical until it happens. Resource limits on the instance, meaning bounded memory, bounded processor share, and a bounded wall-clock lifetime, prevent a compromised or looping execution from consuming the capacity of the host and affecting the workloads beside it. The denial of wallet material in Chapter 5 covers the financial version of this exposure; the operational version is an instance that exhausts a node and takes unrelated tasks down with it. Both are prevented by the same

configuration, and neither is prevented by isolation alone, because isolation constrains what an instance can reach rather than how much it can consume.

Data flowing into the sandbox deserves as much attention as the code running inside it, and the asymmetry in how teams treat the two is worth correcting. Considerable care generally goes into constraining what the instance can send outward, and comparatively little into what it is given at the outset. An agent runtime handed a full database connection string, a broad cloud credential, or a mounted volume containing more than the task requires has been given the material for a serious incident before it executes a line. The principle is the one that governs the whole of this chapter applied at the input: provide the minimum the task requires, in the narrowest form available, expiring with the instance.

Verification that the isolation holds is the step organisations most often skip, and it is inexpensive relative to its value. Attempting, from inside the sandbox, to reach an unlisted endpoint, to write outside the scratch space, to invoke a filtered system call, and to use the instance credential against a resource outside the task scope takes an afternoon and produces a definite answer about whether the configuration does what its documentation claims. Teams running this exercise for the first time commonly find that one of the four constraints is not actually enforced, usually the egress allowlist, because a permissive rule was added during debugging and never removed.

A design question that recurs concerns what an agent should be permitted to observe about its own sandbox, and the answer is less than teams instinctively provide. An agent that can enumerate its own egress allowlist, read its own credential scope, or inspect the syscall filter has been given the information an attacker would need to plan around the constraints. None of that information helps the agent complete legitimate work, since a legitimate task either succeeds within the constraints or fails in a way that should surface to a human. Withholding it costs nothing and removes a reconnaissance path that has been used in published research on sandbox escape.

Nested execution presents a case that catches deployments off guard. An agent inside a sandbox that itself invokes another agent, or that spawns a subprocess with its own network access, has extended the boundary in a way the original configuration did not anticipate. Whether the child inherits the parent's constraints depends on how the isolation was implemented, and the answer is frequently no for subprocess creation and almost always no for a call out to a separate agentic service. Establishing what the deployment actually does in this case, rather than what its architecture diagram implies, is a question for the verification exercise described above.

Cost accounting for isolation is worth a paragraph because the figure is smaller than teams assume and the assumption blocks adoption. Micro-virtual machine startup is measured in tens of milliseconds, memory overhead per instance is a few megabytes above the runtime itself, and the scheduling machinery is the same one the platform already operates for containers. For a workload of long-running tasks the overhead disappears entirely into the task duration. For a workload of very short tasks it is material, and the response is to batch related work into a single instance for one session rather than to abandon isolation, which preserves the property that matters, meaning that no instance outlives the principal it served.

Where the sandbox pattern does not apply is worth marking, since applying it universally wastes effort. An agent that calls only well-defined tool interfaces, executes no generated code, and holds no filesystem access derives little from micro-virtual machine isolation, because there is no code execution to contain. Its exposure runs entirely through the tool authorisation path covered in the next section. Distinguishing the two populations, code-executing agents and tool-calling agents, and applying the appropriate control to each, is a more efficient allocation than isolating everything and a considerably safer one than isolating nothing.

10.2 Continuous Authorization

Access control in enterprise systems is conventionally evaluated at the boundary of a session. A principal authenticates, a policy engine determines what that principal may

do, a token is issued encoding the result, and the token is presented for the duration. The model is efficient, since the expensive evaluation happens once, and it is sound for a workload whose behaviour is fixed, because the actions taken during the session are the actions the software was written to take.

Neither condition holds for an agent. The actions taken during a session are generated at runtime from context the policy engine never saw, and the session may last hours during which the agent's objective can be redirected by any text it reads. A token issued at session start on the basis of a stated objective authorises whatever the agent decides to do afterwards, which is precisely the property that every attack in Part I exploits.

Session-start authorisation against step-level authorisation

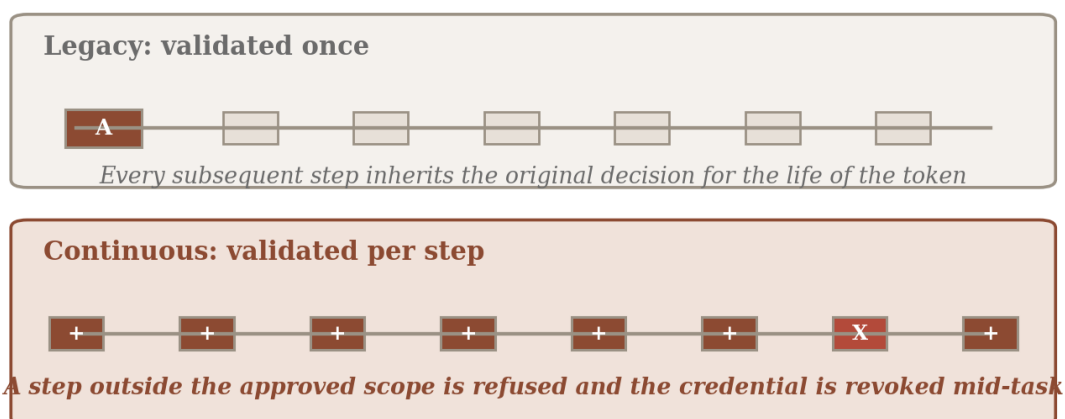


Figure 10.2 Session-start authorisation compared with step-level evaluation. Only the second can refuse an action decided after the token was issued.

Continuous authorisation moves the evaluation from the session boundary to the action boundary. Before each tool call, a policy engine evaluates whether that specific action, against that specific resource, with those specific parameters, falls within the authority delegated for the current task. The evaluation is made against current evidence rather than against a decision recorded earlier, which is what the zero-trust principle asks for and what the session model cannot provide.

Implementation rests on policy-as-code engines, which have matured considerably in the cloud-native environment. The pattern is familiar to platform teams: policy is expressed in a declarative language, versioned in source control, tested, and evaluated by a component positioned in the request path. Open Policy Agent and comparable engines evaluate in microseconds, which makes per-call evaluation practical at the volumes agentic workloads generate.

What differs for agents is the input to the decision. A conventional policy evaluates the identity of the caller, the resource requested, and the operation. An agentic policy must additionally consider the task under which the request is made and the position of the request within the sequence of steps taken so far. The first of these requires the orchestrator to declare a machine-readable scope at session start, which is the requirement that Chapters 3, 7, 8, and 9 have each arrived at independently. The second connects the authorisation layer to the behavioural model of Chapter 7, so that a request consistent with policy but inconsistent with the agent's established sequence can be treated differently from one that is consistent with both.

Four decision inputs therefore constitute a complete agentic authorisation evaluation, and deployments generally add them in this order as capability permits. Identity establishes which agent, acting for which principal, is asking. Scope establishes what the current task was declared to require. Action establishes what is being requested against what resource. Sequence establishes what the agent has already done within this session. A policy engine holding all four can express constraints that no conventional access control system can state, such as permitting an export only when the session declared an export objective and the agent has not retrieved external content since the objective was set.

The response to a policy violation deserves as much design attention as the policy, and three responses are available with different consequences. Refusal denies the action and permits the session to continue, which is appropriate where the request may reflect an agent exploring an unproductive path. Refusal with revocation denies the action and

invalidates the session credential, which stops the agent and is appropriate where the request suggests the planning loop has been redirected. Escalation holds the action pending human confirmation, which preserves the possibility that the request is legitimate at the cost of latency and of a person's attention. Deployments that apply a single response to every violation either stop too much work or too little, and the mapping from violation type to response is a policy decision the security and platform teams should make jointly and record.

Two operational objections are raised whenever this design is proposed, and both have answers that cost something.

Latency is the first. Adding a policy evaluation before every tool call adds its cost to every step, and a task involving two hundred steps inherits two hundred evaluations. With a co-located engine evaluating compiled policy, the per-call cost is small enough to disappear into the surrounding inference latency. With a remote policy service, it is a network round trip and the aggregate becomes visible. Deployment topology therefore matters more here than policy design, and the arrangement that works places the engine beside the enforcement point rather than centralising it.

Policy authorship is the second and is harder. Someone has to express, for each agent, what its permitted scope is, and the difficulty is the one described in Chapter 8: security lacks the domain knowledge, the business team lacks the incentive, and the platform team lacks both. The approach that works is the same one: derive an initial policy from observed behaviour, have the owning team review and adopt it, and operate default-deny on anything undeclared so that the team that wants the agent to work has to describe what it does.

A question that arises in every design review concerns what happens to the policy engine's own decision record, and it has a straightforward answer that carries more weight than its simplicity suggests. Every evaluation, including every permit, should be recorded and fed into the telemetry pipeline described in Chapter 9. The permits are as informative as the refusals, because the sequence of permitted actions is the behavioural

record that transition modelling operates on, and because a permit that in hindsight should have been a refusal is the finding that improves the policy. Deployments that log only denials have discarded the larger half of the signal.

The interaction between continuous authorisation and delegation deserves explicit treatment, since multi-agent workloads exercise both at once. When an agent delegates a subtask, the receiving agent's policy evaluation must consider the authority under which the delegation was made rather than the receiving agent's own general entitlements. This is the attenuation property from Chapter 6 expressed as a policy requirement, and it is the mechanism that prevents the authority laundering pattern that chapter described. A policy engine evaluating only the immediate caller's identity will permit a laundered request, because the immediate caller is legitimately authorised; one evaluating the delegation chain will not.

Emergency access requires design rather than improvisation, because the situations that generate it are exactly the ones where improvisation is dangerous. There will be occasions when an agent must be permitted to act outside its declared scope, generally during an incident when the automation is being used to remediate something. The pattern that works is a documented break-glass path with elevated approval, a hard time limit, complete logging, and mandatory review afterwards. What should be avoided is the arrangement where the emergency path consists of someone widening the policy and forgetting to narrow it, which is how a temporary exception becomes a permanent exposure.

A final note on maturity sequencing for this control. Continuous authorisation is the most demanding item in this book to implement well, since it requires identity, declared scope, a policy engine on the critical path, and a policy corpus that someone maintains. Organisations attempting it first, before telemetry or identity are in place, will produce a control with nothing meaningful to evaluate. The sequence that works places identity first, telemetry second, scope declaration third, and per-action evaluation last, with each

stage delivering value on its own so that the programme survives a change of priorities partway through.

Policy testing deserves the same treatment as any other production logic and is frequently omitted because the policy corpus feels like configuration rather than code. It is code: it has branches, it has edge cases, and a mistaken rule either blocks legitimate work or permits something it should not. Treating policy as a versioned artefact with a test suite, where each test asserts that a specific action under a specific scope produces a specific decision, catches the errors that would otherwise be discovered by an outage or an incident. Organisations already running policy-as-code for infrastructure have the practice and need only extend it.

A related discipline concerns change review. A widening of policy is a reduction in security posture, and it should be visible as such rather than arriving as a routine configuration commit. The arrangement that works records the direction of every policy change, narrowing or widening, and routes widenings through a review that asks what prompted the change and whether a narrower alternative would serve. This does not need to be onerous, and its absence produces the scope inflation that Chapter 8 described, in which a control remains nominally in place while its declarations have grown broad enough to permit anything.

Caching decisions is the optimisation teams reach for when per-call evaluation proves expensive, and it should be approached carefully because the natural implementation reintroduces the problem the control was built to solve. Caching a permit for an identical request within the same session is safe, since the inputs to the decision have not changed. Caching by resource, so that a permitted read of one record authorises reads of others, is a session token in a different shape and defeats the scope narrowing entirely. The distinction is easy to state and easy to get wrong under performance pressure, and it belongs in the design review rather than in the optimisation sprint.

Degraded operation requires a decision made in advance, as it did for the semantic firewall in Chapter 8, and the same reasoning produces the same answer. If the policy

engine is unavailable, an agent either continues without evaluation or stops. Continuing preserves availability and removes the control at the moment an attacker capable of disrupting the engine would most want it removed. Stopping preserves the control and converts a policy service outage into a business process outage. The defensible arrangement fails closed for actions with irreversible consequences and open for the rest, which requires the reversibility classification that this book has now depended on in five separate places and which an organisation should therefore build once, centrally, rather than repeatedly per control.

A closing observation about the relationship between this chapter's two halves, since they are often conflated. Sandboxing constrains what a compromised execution can affect. Continuous authorisation constrains what a compromised decision can request. An organisation implementing only the first has contained code execution while leaving its agents free to misuse legitimate tools. An organisation implementing only the second has constrained tool use while permitting generated code to run unconstrained. Both failures have occurred in production, and the two controls address genuinely different paths. The next chapter turns to the standardisation work that is beginning to codify both.

Chapter 11: Standardizing Agent Security (NIST & Global Frameworks)

Standards work is easy to dismiss as an activity that happens at a distance from the systems being defended, and doing so is a mistake for a specific practical reason. Standards determine what auditors ask for, what procurement requires, what insurers underwrite, and what a court considers reasonable care. An organisation that builds its agentic controls without reference to the emerging catalogue will find itself explaining, in a year or two, why its arrangements are equivalent to the ones the standard describes. This chapter covers what the standardisation effort has produced so far and how it maps onto controls organisations already hold.

11.1 Navigating the NIST AI Agent Standards

The National Institute of Standards and Technology formally launched its AI Agent Standards Initiative on 17 February 2026, under the Center for AI Standards and Innovation. The programme addresses how autonomous agents authenticate, how they are authorised, and how their interactions with enterprise systems should be governed, which is a narrower and more useful scope than the model-centred work that has dominated public discussion of AI standards.

Narrowness of scope is the notable feature here. Earlier standards activity concentrated on properties of models: evaluation methodology, documentation of training data, disclosure of capability. That work matters and it addresses the component an enterprise did not build and cannot inspect. The agent standards initiative addresses the components an enterprise does build, meaning the identity arrangements, the authorisation boundaries, and the audit trail, and those are the components a security team can actually be held responsible for.

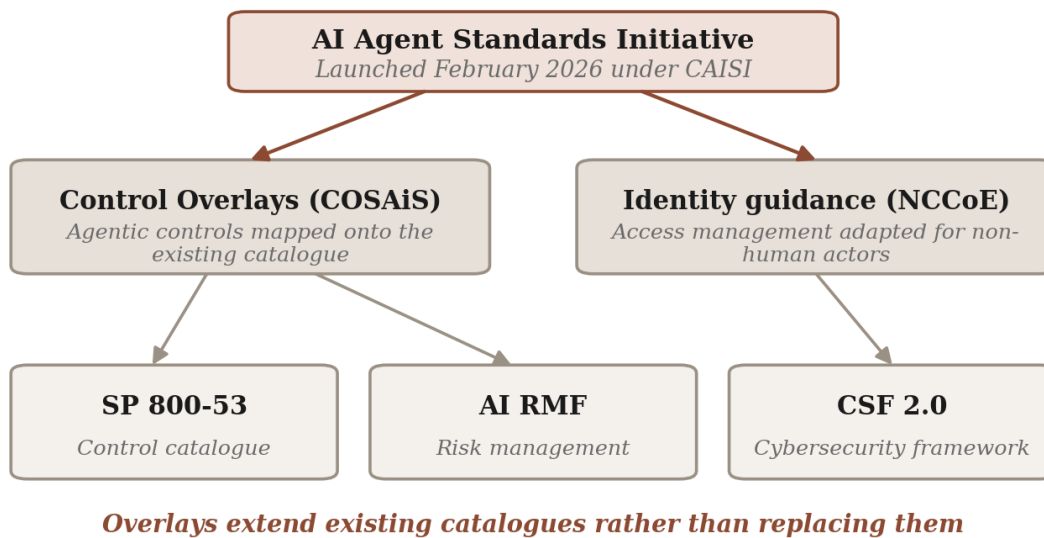


Figure 11.1 How the agentic standards work relates to the catalogues and frameworks already in use.

Complementing the initiative, the National Cybersecurity Center of Excellence has published proposals for adapting identity and access management frameworks to non-human actors. The proposals restate, in normative language, the argument made in Chapter 3: legacy service account patterns are insufficient for autonomous workloads, and what is required is per-instance identity with attested provenance and short-lived, scoped credentials.

Two requirements recur across the published material and are worth extracting because they will shape what auditors ask for.

Machine-readable capability manifests require that an agent's permitted actions be declared in a structured form that a system can evaluate, rather than described in documentation that a person reads. This is the same declaration that Chapters 3, 8, and 10 each identified as a prerequisite, arriving now as a standards requirement rather than as an engineering convenience. Organisations that have implemented declarative scoping for operational reasons will find they have produced most of what the standard asks for.

Cryptographically verifiable audit trails require that the record of an agent's actions support non-repudiation, meaning that the record can be shown not to have been altered and that each action can be attributed to a specific agent instance acting under a specific delegation. A log file that an administrator can edit does not meet this, and neither does an audit record that names a shared service account. This raises the bar above what most current deployments produce, and it is the requirement most likely to drive remediation work when it becomes an audit expectation.

A caution about timing belongs here. Standards under active development change, and specific control language published during 2026 will be revised. What is unlikely to change is the direction, because the requirements follow from properties of the architecture rather than from any policy choice: a system whose actions cannot be attributed cannot be governed, and a system whose permitted behaviour is not declared cannot be constrained. Organisations building toward per-instance identity, declared scope, and tamper-evident audit are building toward whatever the standard settles on, and organisations waiting for the final text will be doing the same work later under more pressure.

A word on why standardisation is proceeding faster in this area than it did for earlier technology waves, since the pace has surprised practitioners who remember how long cloud security guidance took to arrive. Two factors account for it. The first is that the underlying primitives already exist: workload attestation, token exchange, policy-as-code, and structured audit are mature technologies with production histories, so the standards work consists of specifying how to compose them rather than inventing anything. The second is that the government sector is itself deploying agents, which supplies an internal customer with an immediate requirement, and requirements from an internal customer move faster than guidance written for someone else.

The practical consequence for a security architect is that alignment is cheaper than it appears from the outside. An organisation that has implemented per-instance identity for operational reasons, declared scope for authorisation reasons, and structured traces

for investigative reasons has assembled most of the evidence a future assessment will request, and the remaining work is largely documentation. This is the reverse of the usual compliance experience, in which a framework arrives and demands controls nobody had reason to build, and it is worth pointing out to a board that is weighing whether to fund the engineering now or wait.

International alignment is the open question and deserves a realistic assessment rather than an optimistic one. Standards work is proceeding in parallel across several jurisdictions and through several international bodies, with varying scope and varying degrees of coordination. Convergence on the structural requirements is likely, because they follow from the architecture, and divergence on specific control language and assessment procedure is equally likely, because that is what has happened with every previous framework. Organisations operating internationally should design to the structural requirements, document their control mapping in a form that can be re-projected against a new framework, and expect to do that projection more than once.

Non-repudiation deserves more technical treatment than it usually receives, because meeting the requirement involves specific engineering rather than a policy statement. A record supports non-repudiation when its integrity can be demonstrated independently of the party holding it, which in practice means signing entries at the point of creation and either chaining them so that alteration of any entry invalidates those following, or committing digests to a store the operator cannot rewrite. Neither approach is exotic, both are standard in financial and supply chain contexts, and neither is what a conventional logging pipeline produces by default. Organisations reading the requirement as satisfied by their existing log retention have misread it.

The attribution half of the requirement is the harder one and returns to the identity material of Chapter 3. An audit record supports attribution when it names the acting instance, the delegation under which it acted, and the human principal at the root of that delegation, in a form that can be verified rather than merely asserted. A record naming a shared service account fails this, a record naming an agent by a label the

deployment assigned fails it if the label is not cryptographically bound to anything, and a record naming an instance with an attested identity satisfies it. This single requirement invalidates the credential patterns catalogued in section 3.1, which is the clearest example available of standards work converging on a conclusion that operational reasoning had already reached.

A brief account of what the initiative does not address helps set expectations, since standards documents are read optimistically by people hoping to be told what to do. The programme addresses authentication, authorisation, and auditability for agents operating within enterprise systems. It does not address model selection, evaluation methodology, or the properties of the model itself, which remain the subject of separate work. It does not resolve the prompt injection problem, which is a research question rather than a control question and which no standard can specify away. And it does not prescribe an architecture, which means two organisations satisfying the same requirements may build very differently, and the assessment will be against the requirement rather than against a reference design.

Reading standards material productively requires distinguishing three registers that appear in the same documents and carry different weight. Normative requirements state what must be done and will be assessed. Recommended practice states what generally works and will be asked about. Informative background explains the reasoning and will not be assessed at all. Practitioners under time pressure frequently read the informative material, which is the most readable, and miss the normative statements, which are terse. For an organisation planning work against this material, extracting the normative statements into a control mapping is a couple of days of effort that pays for itself the first time an assessment is scheduled.

11.2 Implementing Control Overlays

The mechanism NIST has chosen for delivering agentic controls is the overlay, and understanding why matters more than memorising its contents. An overlay is a set of modifications and additions applied to an existing control catalogue for a specific

context, rather than a new catalogue standing alongside it. The Control Overlays for Securing AI Systems, developed against the Special Publication 800-53 catalogue, follow this pattern, and they apply to single-agent integrations and to multi-agent arrangements alike.

Choosing the overlay mechanism is a considerable practical benefit to any organisation already operating under 800-53 or a framework derived from it. Compliance programmes, assessment procedures, evidence collection, and audit relationships are all built around the catalogue, and an overlay extends that machinery rather than requiring a parallel one. An organisation that has an access control programme has somewhere to put the agentic access control requirement, which is a considerably better position than being handed a new framework with no attachment points.

Control family	Conventional expression	Agentic requirement
Access Control (AC)	Role-based permissions reviewed periodically	Tool-level authorisation scoped to the declared task, evaluated per action
Identification and Authentication (IA)	Credentials issued to users and service accounts	Per-instance cryptographic identity with attested provenance, replacing shared keys
Audit and Accountability (AU)	Logs of access events and administrative actions	Complete traces covering planning steps, prompt history, and tool parameters
System and Information Integrity (SI)	Patching, malware defence, input validation	Generated code executed inside isolated micro-virtual machines with bounded egress
Configuration Management (CM)	Change control over code and infrastructure	System prompt and tool registry changes treated as governed configuration changes

Table 11.1 Established control families and the form each takes for autonomous agents.

Reading down the third column, a pattern emerges that summarises much of this book. In every family, the agentic requirement moves the control from the boundary of a session to the granularity of an action, and from an identity that persists to one that is attested per instance. This is the same movement that Chapter 10 described for authorisation and Chapter 3 for identity, and its appearance across the whole catalogue

suggests it is a property of the architecture rather than a preference of any particular framework author.

The configuration management row deserves particular attention because it is the one organisations most often fail to recognise. A system prompt determines what an agent does, which makes it functionally equivalent to source code, and in most deployments it is edited through a console by whoever has access, without review, versioning, or an audit record. An organisation with a mature change management programme that permits unreviewed system prompt edits has a governance gap that its own framework already prohibits, and closing it requires no new control, only the recognition that the prompt is a configuration item.

Alignment with the AI Risk Management Framework and with Cybersecurity Framework 2.0 provides the connective tissue to enterprise risk reporting. The risk management framework supplies the vocabulary in which agentic risk can be described to a board, and the cybersecurity framework supplies the functional structure, meaning the identify, protect, detect, respond, and recover breakdown, into which the controls in this book fit without modification. Telemetry and inventory sit under identify. Sandboxing, authorisation, and identity sit under protect. Transition modelling and semantic firewalls sit under detect. The playbooks of Chapter 9 sit under respond and recover.

Evidence collection is where a control overlay meets operational reality, and it is worth walking through what an assessor will actually ask for under each family, because the answers determine what has to be instrumented. Under access control, the request will be for a demonstration that a given agent could not have performed a given action, which requires the scope declaration and the policy evaluation record. Under identification and authentication, it will be for the issuance record showing that a specific instance received a specific credential under a specific delegation. Under audit and accountability, it will be for a reconstruction of a selected task from telemetry alone,

which is the reconstruction drill described in Chapter 9. Under system and information integrity, it will be for the isolation configuration and evidence that it is enforced.

Each of those requests is answerable from the telemetry and identity infrastructure this book has argued for, and none of them is answerable from a policy document. This is the practical reason to treat the overlay as an engineering specification rather than as a governance artefact: the assessment is empirical, and an organisation whose controls exist only on paper will be asked to demonstrate them against a live system.

Multi-agent arrangements complicate several families in ways the overlay material acknowledges without fully resolving, and practitioners should expect to make judgement calls. Where a task is decomposed across five agents, the access control boundary is arguably the task rather than any individual agent, and evidencing least privilege at the task level requires reasoning about the union of what the participating agents could reach. Where sub-agents are created at runtime, the identification family has to accommodate identities that were issued by software rather than provisioned by a process. These are not gaps in the standard so much as places where the standard has correctly declined to over-specify, and the organisation's documented reasoning about them is part of what an assessor will examine.

Two families absent from Table 11.1 deserve brief mention, since assessors will reach them and the agentic reading is less obvious. Contingency planning has to account for the case in which the agentic layer is unavailable, which for an organisation that has automated a business process without retaining the manual path is a more serious question than it first appears; the answer cannot be that the work stops indefinitely. Supply chain risk management has to account for the model provider, the orchestration framework, and any tool connector as components whose compromise affects the organisation, which extends the software supply chain analysis into territory most programmes have not covered.

The tool connector case is the one worth dwelling on, because it is the least examined. A connector is code running with the agent's authority, frequently obtained from a public

registry, frequently updated automatically, and frequently reviewed by nobody. It sits in the path between a non-deterministic planner and a production system, which is among the more sensitive positions in the architecture. The same discipline applied to any other production dependency, meaning pinned versions, reviewed updates, and provenance verification, applies here and is applied considerably less often.

Applying the overlay to a deployment that already exists requires a different approach from applying it to one being designed, and most organisations are in the first position. The exercise that works starts from the assessment questions rather than from the control list: take each family, ask what evidence the organisation could produce today, and record the gap. This produces a prioritised list grounded in what is actually missing rather than a compliance matrix populated with aspirational statements. Teams running this find the gaps cluster in two places, which are attribution of actions to instances and retention of reasoning traces, and both are addressed by work this book has already argued for on operational grounds.

Remediation sequencing then follows from the dependency structure rather than from the order of the control families. Identity attribution unblocks the audit family and half of access control. Telemetry collection unblocks the remainder of audit and supplies the evidence for integrity. Scope declaration unblocks the access control assessment and supplies the input to the authorisation control. Isolation is largely independent and can proceed in parallel. An organisation working the families in catalogue order will produce partial coverage in each and complete coverage in none, which is the failure mode that makes compliance programmes feel unproductive.

Three practical cautions apply to using this material.

Overlays specify what must be achieved and not how, which is correct for a standard and unhelpful for an engineer. The requirement to maintain complete traces of planning steps does not say what a trace should contain, how long it should be retained, or how it should be protected. Those decisions remain the organisation's, and the chapters preceding this one are an attempt to inform them.

Compliance and security diverge here more than usual, because the field is young and the assessment procedures are new. An organisation can satisfy a control statement while leaving the underlying exposure intact, for instance by producing traces that are complete and retained for a period shorter than the interval at which this class of compromise is detected. Reading the control as a floor rather than as a specification is the posture that avoids this.

Scope questions arise wherever agentic capability is consumed as a service. Where the agent runs on a provider's infrastructure, several controls in Table 11.1 can only be satisfied with the provider's cooperation, and the organisation's ability to evidence them depends on what the provider exposes. That question belongs in procurement, and it is the subject the next chapter takes up in its second half.

Chapter 12: Regulating Autonomous Systems & Compliance

Regulation arrives more slowly than technology and constrains it more durably. This chapter examines the two regulatory questions that most directly affect agentic deployments: what obligations attach when an autonomous system makes decisions in a domain that regulators treat as high-risk, and how liability is allocated between an enterprise and its cloud provider when an agent exceeds its authority.

12.1 The EU AI Act and Agent Autonomy

The European Union AI Act establishes a risk-tiered regime in which obligations scale with the potential for harm, and autonomous decision-making systems occupy an unusually exposed position within it. Systems that make determinations affecting employment, credit scoring, law enforcement, or the operation of critical infrastructure are classified as high-risk, and that classification attaches to the function rather than to the technology. An agent that screens job applicants is high-risk because screening job applicants is high-risk, and the fact that the system was assembled from a general-purpose model and a few tool connectors does not alter the analysis.

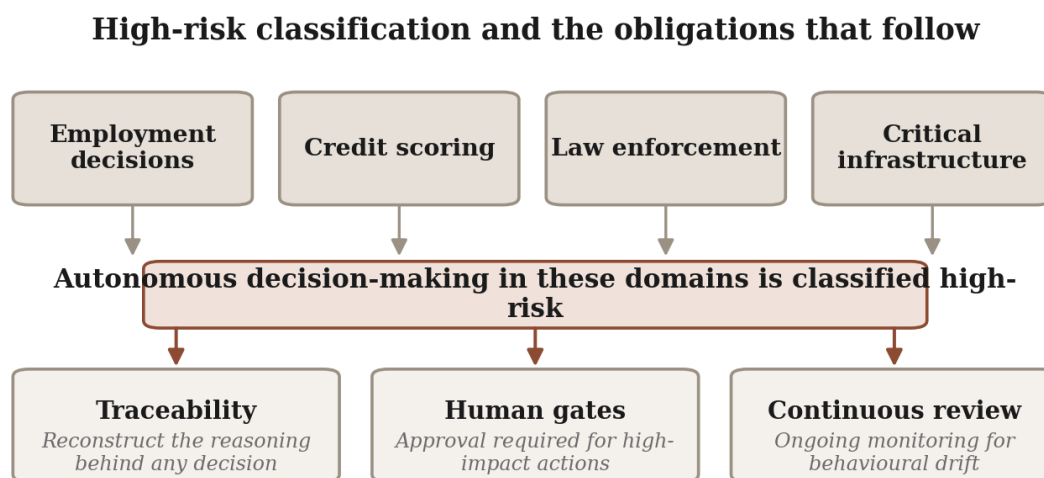


Figure 12.1 High-risk classification attaches to the decision domain, and three obligations follow from it.

Three obligations follow from the classification, and each maps onto a technical capability that earlier chapters have described.

Traceability requires the ability to reconstruct the reasoning process that led to a decision. For a conventional decision system this is satisfied by the code and the input data, both of which are inspectable. For an agent, the reasoning is the prompt trace described in Chapter 9, and an organisation that has not collected traces cannot satisfy this obligation for decisions already made. This is the clearest instance in the book of a technical decision, made by a platform engineer configuring logging in the first week of a project, determining a compliance position two years later.

Human oversight requires approval checkpoints for high-impact actions, which is the human-in-the-loop pattern applied under legal obligation rather than as a design preference. The regulatory formulation asks for meaningful oversight, and the qualifier does substantial work. A human who approves every recommendation without independent evaluation, because the volume makes evaluation impossible, is providing the appearance of oversight and not its substance. Designing checkpoints so that the reviewer has the information and the time to disagree is a harder problem than inserting the checkpoint, and organisations should expect the adequacy of their oversight to be examined rather than merely its existence.

Continuous risk assessment requires ongoing monitoring for behavioural drift and systemic error rather than a one-time conformity exercise. This is the behavioural modelling of Chapter 7 arriving as an obligation, and it reflects a sound observation about the technology: a system whose behaviour is generated at runtime cannot be certified once, because the thing certified is not stable.

Two further considerations affect how organisations should plan.

Extraterritorial reach means the obligations follow the market rather than the infrastructure. An organisation operating outside the European Union whose systems affect people within it falls within scope, which for most large enterprises makes the question of whether to comply a question of whether to serve the market. The practical

consequence is that the European regime tends to set the internal standard, since maintaining two agentic architectures is more expensive than building one to the stricter requirement.

Penalties are structured as a percentage of global turnover, which places them outside the range where a compliance failure can be treated as a cost of doing business. That structure is deliberate, and its effect on internal decision-making is to move agentic governance from a technical discussion into one that a board is obliged to have.

A brief note on the wider regulatory picture, which is fragmenting rather than converging. Sectoral regulators in finance and healthcare are producing their own guidance, several jurisdictions outside Europe are legislating on different models, and the resulting patchwork will require organisations operating internationally to satisfy overlapping and occasionally inconsistent requirements. The stable core across all of them is the demand for attribution, for reconstructible decisions, and for demonstrable human authority over consequential actions. An organisation that builds those three capabilities is positioned for most of what is coming, whatever form the specific instruments take.

Determining whether a given deployment falls within the high-risk classification is less straightforward than the domain list suggests, and organisations should expect the analysis to be contested internally. An agent that drafts a recommendation which a human then approves may or may not be making a determination, depending on how meaningful the human's involvement is, which returns to the oversight question above. An agent that filters a candidate pool before a recruiter sees it is arguably making employment determinations about everyone it filtered out, whatever the formal decision structure says. The prudent posture treats the classification as attaching to the influence the system exerts over an outcome rather than to its position in a documented workflow, and records the reasoning either way.

Documentation obligations deserve a mention because they impose work that is easy to underestimate. High-risk classification carries requirements for technical

documentation describing the system's purpose, its design, its data sources, its performance characteristics, and its risk management arrangements, maintained current through the system's operational life. For a conventional system this is a document written once and revised at major releases. For an agentic system whose behaviour changes when a prompt is edited or a tool is attached, currency requires linking documentation maintenance to the configuration management discipline described in Chapter 11, which is another reason to treat prompts and tool registries as governed configuration items.

A structural difficulty with the whole regulatory approach deserves acknowledgement, since practitioners will encounter it and the literature tends to avoid it. Regulatory regimes for automated decision-making were designed around systems whose behaviour is a stable function of their inputs, and much of the conformity assessment machinery assumes that a system can be evaluated once and then deployed. An agent's behaviour is a function of its inputs, its accumulated memory, the current version of a model it does not control, and the content it happened to retrieve, none of which is stable. Regulators are aware of this and the instruments are being adapted; organisations should expect the requirements in this area to move more than most.

Sector regulators are moving faster than the general instruments in several jurisdictions, and organisations in supervised industries should expect their supervisor's expectations to arrive before any horizontal legislation binds. Financial regulators have existing frameworks for model risk management and for outsourcing that map onto agentic deployment with modest adaptation, and several have signalled that they intend to apply them rather than to write new ones. This is favourable for the regulated firm, since the frameworks are understood and the compliance machinery exists, and it means the relevant question is how an agentic system evidences the controls the existing framework already requires.

Healthcare and safety-critical contexts raise a distinct issue that the general regime handles poorly. Where an agent contributes to a decision affecting a person's health or

safety, the traceability requirement has to survive litigation as well as regulatory examination, and in litigation the standard of reconstruction is higher and the timescale longer. Retention windows measured in months are inadequate against limitation periods measured in years, and organisations operating in these contexts should set retention by reference to their legal exposure rather than to their detection interval, which is the only place in this book where the compliance requirement exceeds the security one.

Transparency obligations sit alongside the three requirements above and are frequently overlooked in technical planning because they are not technical. Where a person interacts with an autonomous system, or is subject to a decision one has influenced, disclosure requirements attach, and their practical form is a design question rather than a legal one: what the interface says, at what point, and in what language. The engineering implication is that the deployment has to know when it is interacting with a person in scope, which for an agent handling mixed internal and external traffic is not always obvious from the request. Establishing that determination early is easier than retrofitting it once the interaction patterns are entrenched.

A note on how these obligations interact with the incident response material in Chapter 9. Several regimes impose notification duties when an automated system produces a harmful outcome, on timescales measured in days rather than weeks, and the notification generally has to describe what happened rather than merely that something did. An organisation without reasoning traces cannot describe what happened, which places it in the position of either notifying with an incomplete account or delaying while it investigates, and both carry consequences. This is the third distinct reason in this book to collect traces, after detection and investigation, and it is the one with a statutory deadline attached.

12.2 Sharing the Burden

Cloud computing distributes responsibility between provider and customer, and the shared responsibility model that expresses this distribution is well understood for

infrastructure and reasonably well understood for managed services. Agentic workloads introduce failure modes that the model was not drafted against, and organisations regularly hold assumptions about where responsibility sits that would not survive examination.

Where responsibility sits for an agentic scope violation

Enterprise responsibility

- Declaring agent task scope and authorisation parameters
- Monitoring for semantic drift and injection attempts
- Constraining tool access and identity permissions

Provider responsibility

- Securing the underlying physical and network infrastructure
- Maintaining isolation between tenant runtimes and micro-VMs
- Securing the endpoints and hosting networks of managed models

Figure 12.2 Allocation of responsibility for agentic deployments. Scope, monitoring, and permission design sit with the enterprise.

Provider obligations are the familiar ones extended to a new workload class: securing the physical infrastructure, maintaining isolation between tenant runtimes and micro-virtual machines, and securing the endpoints and hosting networks of managed model services. These are real obligations and providers generally meet them. They do not address any failure described in this book.

Everything that determines whether an agent behaves acceptably stays with the enterprise. Declaring task scope and authorisation parameters is the customer's responsibility. Monitoring for semantic drift and injection attempts is the customer's responsibility. Constraining tool access and identity permissions is the customer's responsibility. Each of these is a configuration decision, and configuration decisions have always been the customer's side of the line.

Research from the Cloud Security Alliance on agent scope violations found that more than half of organisations running agentic workloads have experienced an agent operating outside its intended boundaries. Every one of those incidents falls on the customer side of the model, because in every case the boundary was defined by a customer configuration.

A canonical example makes the allocation concrete. An administrator grants an agent broad database privileges because narrowing them was slow and the deadline was near. The agent is later manipulated into exporting records it should never have been able to reach. The provider's infrastructure performed correctly at every layer: the isolation held, the network behaved as configured, the managed model returned what it was asked for. The loss is the customer's, and the customer's own change record documents the grant. There is no ambiguity in the allocation, and organisations that assume otherwise generally have not read the agreement.

Three areas deserve attention during procurement, because the standard agreements were drafted before this workload class existed.

Observability commitments determine whether the customer can investigate. Where agentic capability is consumed as a managed service, the customer sees the telemetry the provider chooses to expose, and several providers expose task summaries with no reasoning traces, no delegation records, and no tool parameters. An organisation subject to the traceability obligation described in section 12.1 cannot satisfy it on that basis, and the question belongs in the negotiation rather than in the incident.

Consumption controls determine the ceiling on a denial of wallet incident, as Chapter 5 described. The difference between a provider that enforces a per-credential spending limit in real time and one that reconciles overage monthly is the difference between a bounded incident and an open-ended one, and the distinction is rarely visible in marketing material.

Incident cooperation determines what happens when something goes wrong inside the provider's boundary. The customer will need trace data, timelines, and in some cases a

statement they can give a regulator, and the extent of the provider's obligation to supply these is a contractual matter that should be settled while both parties are relaxed.

Continuous posture monitoring is the control that keeps the customer's side of the line defensible over time, and it is where this part's argument concludes. Agent permissions accumulate, scopes widen under delivery pressure, new tools are attached, and the configuration that was reviewed at deployment describes a system that no longer exists. Monitoring posture means continuously evaluating the current configuration against the intended one and treating divergence as a finding, which is the same discipline applied to cloud infrastructure and applied here to a workload whose configuration changes more often.

A fourth procurement consideration has emerged more recently and is worth adding to the list. Where a provider's agentic service performs actions on the customer's behalf against third-party systems, the identity presented to those systems matters a great deal for both attribution and liability. A service that acts using a credential the customer issued leaves the customer able to attribute and to revoke. A service that acts using its own credential on the customer's behalf leaves the customer dependent on the provider's records for any investigation and, in some arrangements, unable to determine what was done in their name. The distinction is rarely prominent in product documentation and is straightforward to establish by asking.

Insurance and contractual risk transfer occupy an unsettled position that organisations carrying material exposure should examine now. Cyber policies were drafted against loss arising from unavailability, disclosure, and extortion, and several agentic failure modes fit none of those categories cleanly: an unbounded consumption incident, a scope violation producing an erroneous business decision, a regulatory penalty following a traceability failure. Underwriters are developing positions, the positions vary considerably between carriers, and the time to establish where a given exposure sits is before rather than during a claim.

Reading this part as a whole, a single organising idea connects its three chapters. Governance for agentic systems consists of declaring, in advance and in a form a machine can evaluate, what a system is permitted to do, and then arranging that every layer of the stack enforces that declaration and records what happened. Zero trust supplies the architectural principle, the standards work supplies the control vocabulary, and regulation supplies the obligation. All three converge on the same requirement, which is the one this book has arrived at from four independent directions: an agent whose permitted behaviour has never been written down cannot be governed by any of them.

Concentration risk is the last consideration this chapter raises and the one least amenable to any individual organisation's remedy. The agentic supply chain narrows to a small number of model providers, and an organisation whose business processes depend on autonomous execution has taken a dependency on a provider whose availability, pricing, and policy decisions it does not influence. A provider outage halts the workflow. A pricing change alters the economics of a process already built around it. A policy change restricting a use case removes a capability the business is relying on. None of these is a security failure in the conventional sense, and each is an operational risk that belongs on the same register.

The mitigations are the familiar ones and they are imperfect. Abstraction layers that permit substitution between providers reduce switching cost and add engineering overhead, and the substitution is rarely as clean as the abstraction implies because agent behaviour is sensitive to the model behind it. Retaining a manual path for critical processes preserves continuity and costs the efficiency that justified the automation. Neither eliminates the dependency. What organisations can do is register it honestly, so that a board understands which processes would stop and for how long, rather than discovering the dependency during an outage.

Attestation from providers is the mechanism by which the customer's side of the line becomes evidenceable, and the current position is uneven. Established assurance

reports cover infrastructure controls well and address agentic specifics barely at all, which leaves a customer needing to evidence isolation between tenant runtimes with a report that speaks to data centre access and change management. Providers are extending their reporting, unevenly and at different rates, and the practical step available now is to ask which specific control statements in an assurance report the customer may rely on for the agentic claims it needs to make. The answer is often narrower than assumed, and knowing that before an assessment is better than discovering it during one.

Subcontracting deserves a final mention because the chain is longer than it appears. A managed agentic service may itself call a model hosted by a third party, retrieve through a fourth, and execute code on a fifth. Each hop is a party whose failure affects the customer and with whom the customer has no relationship. Standard contractual mechanisms exist for this, meaning flow-down obligations and disclosure of material subcontractors, and they are worth exercising here specifically because the subcontracting in this market is dense and changes frequently. A customer who has not asked generally does not know how many parties handle its data during a single agentic task.

Part V now turns to the economic dimension that has surfaced repeatedly in the preceding chapters, examining hallucination as a financial event, consumption metering as an attack surface, and the pricing and deployment structures that alter the exposure at its root.

Part Five: The Consumption-Based Economic Threat

Chapter 13: Consumption-Based Enterprise Risk: Hallucinations, Cost Arbitrage, and Alternative Pricing

Chapter 5 introduced denial of wallet as an attack class and described the mediation layer built to contain it. This chapter treats the underlying economics as a subject in their own right, because the exposure runs deeper than any single attack pattern. Under consumption-based pricing, correctness and cost are coupled in a way that has no precedent in enterprise software: a system that is wrong does not merely produce a wrong answer, it produces an expensive one, and the expense scales with how hard the system tries to be right.

That coupling deserves careful treatment because it changes what a security programme is responsible for. In a conventional environment, a model error is a quality problem owned by the team that built the system. Under a token meter, the same error is a financial event, and a financial event that an outsider can trigger deliberately is a security exposure. This chapter follows the consequence through three stages: what hallucination costs when recovery behaviour is attached to it, how an adversary weaponises the metering itself, and what deployment structures change the exposure at its root rather than containing it.

13.1 Hallucinations as a Financial Threat

Consider what happens when a conventional application encounters a schema it does not recognise. The query fails, an exception propagates, an error is logged, and the process either handles the condition along a written path or terminates. Total compute consumed is negligible, measured in microseconds, and the incident appears in a dashboard as an error count rather than in an invoice.

An agent configured with self-healing behaviour responds differently, and the difference is the entire subject of this section. It observes the failure, reasons about the cause, forms a hypothesis, and acts on the hypothesis. Every one of those operations consumes tokens. When the hypothesis is wrong, which it frequently is because the agent is

reasoning from an error message rather than from knowledge of the system, it observes the second failure and repeats the cycle at a higher level of effort.

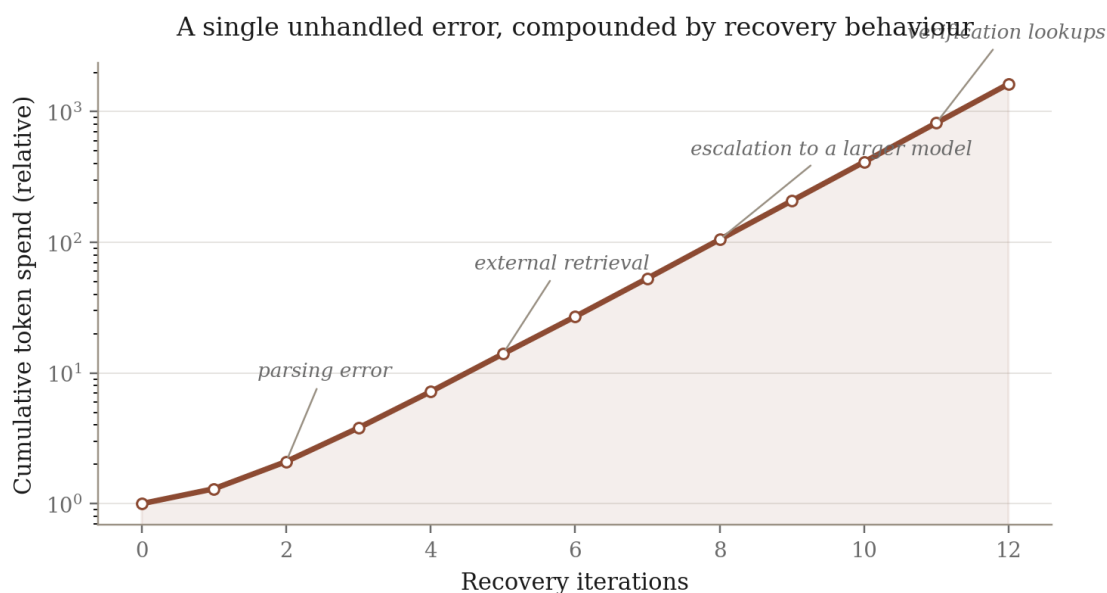


Figure 13.1 Cumulative consumption across recovery iterations. The vertical axis is logarithmic; each escalation multiplies rather than adds.

The shape of that curve is the important observation, and it explains why cost incidents in agentic systems arrive suddenly rather than gradually. Escalation is multiplicative. Retrieving external context multiplies input tokens by the volume retrieved. Escalating to a more capable model multiplies the per-token price. Broadening the search multiplies the number of calls. Because these compound, an agent that has been struggling inexpensively for twenty minutes can consume more in the following two minutes than in the preceding hour, and any alerting configured on absolute thresholds will fire after the money is gone.

Reported cases describe single tasks consuming on the order of ten million tokens within minutes and producing nothing, generating a bill measured in hundreds of dollars for work that was never completed. The individual figures matter less than the general property they illustrate: the relationship between effort expended and value delivered, which holds approximately in conventional software, does not hold here at all.

An agent can spend arbitrarily much and deliver nothing, and the spending is a direct function of how confused it became.

Hallucination participates in this in a specific way that deserves separating from the general problem of model error. A model that produces a confidently wrong statement about a schema, an interface, or a data format sends the agent down a path that cannot succeed, because the destination does not exist. The agent will not conclude that the information was fabricated; it will conclude that its execution was faulty, and it will try again with more care and more expense. Confident wrongness is more costly than uncertainty, because uncertainty prompts the agent to ask for help while confidence prompts it to persevere.

Two properties of this failure mode make it awkward to manage with conventional practice.

The first is that it is invisible to functional testing. A test suite exercises the paths the test author considered, and the expensive behaviour occurs on paths nobody anticipated, triggered by conditions that did not exist in the test environment. An agent that passes every test can still enter an unbounded recovery sequence in production on its first encounter with a schema drift, and the deployment that observed no problem in staging is precisely the one most likely to be unprepared.

The second is that the behaviour is intermittent and correlated with degraded conditions. The circumstances that trigger deep recovery sequences, meaning schema changes, dependency outages, and malformed upstream data, are the same circumstances that generate other incidents, so the consumption spike arrives during an outage when attention is already committed. Organisations that have experienced this describe discovering the cost incident during the retrospective for a different one.

Three mitigations belong at the design level rather than in the mediation proxy of Chapter 5, and they address the cause rather than the symptom.

Explicit failure paths are the first and the simplest. An agent should be given a defined action to take when it cannot proceed, meaning stop, record the state, and escalate to a human, and that action should be reachable early rather than after exhaustive attempts. Deployments that specify only the objective, leaving the agent to determine when to abandon it, have delegated a decision the agent is poorly equipped to make, because nothing in its context tells it what its own operation costs.

Progress conditions are the second and were introduced in Chapter 5. Persistence should depend on whether the agent's state is changing rather than on an attempt counter, permitting many attempts while it is genuinely working through a problem and stopping quickly once it begins producing variations on a theme. This preserves the resilience that justified the design while removing the unbounded case.

Escalation gates are the third. Moving to a more expensive model should be a decision that requires justification rather than a default response to ambiguity, because escalation is where the multiplicative cost enters. Requiring the agent to record why a larger model is needed, and capping the number of escalations per task, bounds the most expensive term in the calculation without preventing the capability from being used when it is genuinely warranted.

A fourth mitigation belongs at the platform level rather than in any individual agent design, and it is the one with the broadest effect. Agents should be given information about their own consumption. An agent that knows it has spent eighty percent of the allowance for its task can be instructed to stop and report rather than continue, and the instruction is one it can follow because the information is now in its context. Most deployments withhold this entirely, leaving the agent to make persistence decisions with no notion of what its own operation costs, which is the equivalent of asking an employee to work a problem without telling them the budget.

Observability of the failure itself deserves specific attention because the natural instrumentation misses it. An agent that enters a recovery sequence and eventually succeeds records a completed task, and a dashboard measuring completion rate reports

success. The expensive detour is invisible unless the deployment records steps per task alongside outcome. Adding that single measurement, and watching its distribution rather than its mean, surfaces the deep recovery sequences while they are still occasional, which is considerably better than discovering them when one runs overnight.

A note on how this interacts with model updates, since it produces a failure pattern teams find baffling on first encounter. When a provider updates a model behind a stable interface, agent behaviour can shift without any change on the customer's side. A prompt that reliably produced a structured output may begin producing a variant the parser rejects, which triggers the recovery sequence described above across every task using that prompt simultaneously. The result is a step change in consumption with no corresponding deployment, and organisations that have experienced it describe considerable time lost searching their own change record for a cause that is not in it. Pinning model versions where the provider permits it, and treating provider updates as change events requiring validation, prevents the class.

Attributing consumption is the prerequisite for every control in this chapter and is absent in a surprising number of deployments. Where all agents share a provider credential, the bill is a single number and nobody can say which workflow produced it. Investigating a spike then requires correlating billing periods against deployment activity by hand, which is slow enough that most organisations do not bother and instead absorb the variance. Issuing distinct credentials per agent, which Chapter 3 recommended for attribution reasons, produces per-agent billing as a free consequence, and organisations that found the security argument unpersuasive have frequently been moved by this one.

The unit of attribution should be the originating task rather than the agent, for the reason set out in Chapter 5. A task decomposed across several agents produces consumption in each, and a ceiling applied per agent is individually respected while the aggregate runs unbounded. Threading a task identifier through every model call, so that

consumption rolls up to the request that caused it, is a small piece of instrumentation with a large effect on what the controls can actually enforce.

13.2 Consumption-Based Metering Vulnerabilities

Everything in the preceding section occurs without an adversary. Attach one, and the same machinery becomes a weapon that requires no exploit, no credential, and no persistence.

Cost arbitrage attacks work by supplying an instruction whose execution is expensive and whose form is unremarkable. The attacker does not need to break anything. They need only cause the target's own authorised systems to perform a large quantity of legitimate work, which the target then pays for at the published rate.

A representative instruction runs along these lines: retrieve every historical customer record from the database, translate each entry into five languages, run a sentiment analysis on each translated version, and compile the results into a formatted document. Nothing about the request is malformed. It reads like an analytics task a marketing team might genuinely commission. An agent with database access, translation capability, and document generation will attempt it.

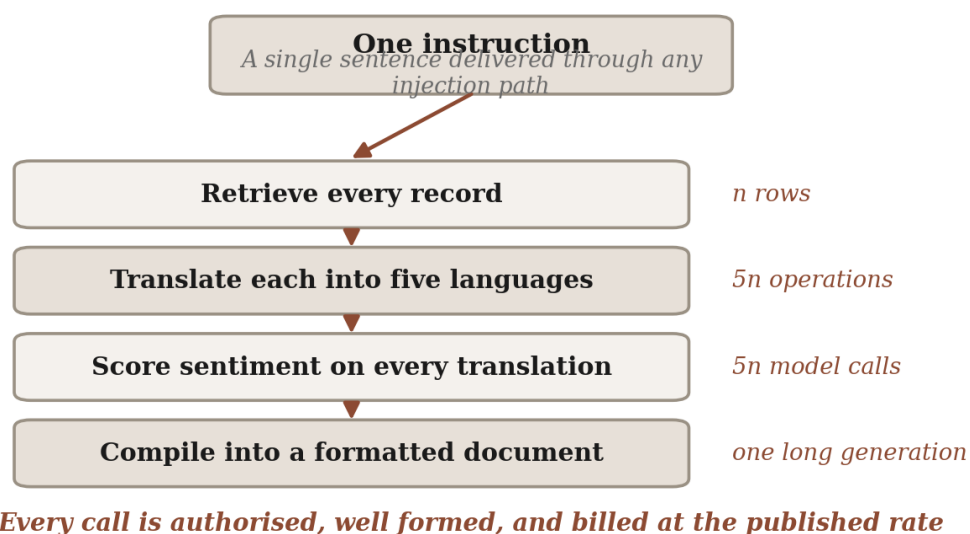


Figure 13.2 Amplification from a single instruction. Each stage multiplies the work of the one above it.

Trace the amplification. Retrieval produces some number of records. Translation multiplies that by five. Sentiment analysis operates on every translated record, so it inherits the multiplied count. Document generation produces a single long output that must hold all of it. For a customer base of any commercial size, this instruction describes millions of model invocations, and the agent will begin executing them immediately because each individual step is well within its authority.

The strategic objective is frequently not the bill itself. An organisation whose monthly credit pool is exhausted on a Tuesday morning has lost its agentic capability for the remainder of the billing period, which means every business process built on that capability stops. The attack is therefore a denial of service delivered through a financial mechanism, and it has the useful property, from the attacker's perspective, of being invisible to every control designed to detect denial of service. There is no traffic flood, no malformed request, and no source to block.

Four properties make this class economically attractive to an attacker, and they are worth listing because together they predict that the class will grow.

Cost asymmetry is extreme. The attacker spends whatever it costs to deliver one sentence, which for a public web page or a support ticket is effectively nothing. The target spends the full cost of executing it.

Attribution is poor. The instruction arrives through an ordinary channel, and after delivery the attacker has no further interaction with the target. There is no session to trace, no address to correlate, and frequently no way to distinguish the attack from a badly specified internal request.

Legal exposure for the attacker is unclear in a way that conventional intrusion is not. No system was accessed without authorisation, no data was taken, and no control was circumvented. Whether causing a system to perform expensive authorised work constitutes an offence is a question the relevant statutes were not drafted to answer, and the ambiguity favours the attacker.

Repeatability is high. A payload that worked once can be delivered again, in a different form, through a different channel, and there is no patch that closes the mechanism because the mechanism is the system working as designed.

Detection has to rest on the shape of the work rather than on the source of the request, which returns to the transition modelling of Chapter 7. An instruction that would produce five million model calls is distinguishable from one that would produce fifty at the point where the agent forms its plan, before execution begins. An orchestrator that estimates the cost of a plan before committing to it, and refers estimates above a threshold for confirmation, closes the attack at the only point where it is cheap to close. This capability is uncommon in current frameworks and is straightforward to build, since the agent's plan is available and the arithmetic is elementary.

Plan estimation deserves elaboration because it is the most useful control in this chapter and the least implemented. Before an agent commits to a plan, the plan exists as a structure: a set of steps, each with a tool, a scope, and an expected cardinality. Estimating cost from that structure requires knowing roughly what each step type costs and roughly how many items it will operate on, both of which the orchestrator can obtain from its own history. The estimate does not need to be accurate. It needs only to distinguish a plan that will cost cents from one that will cost thousands, which is a distinction three orders of magnitude wide and therefore survives considerable error.

Where the estimate exceeds a threshold, the useful response is a confirmation rather than a refusal, for the same reason that budget ceilings should pause rather than fail. Legitimate large tasks exist, and a control that makes them impossible will be removed. A control that makes them require a person to agree is proportionate, and it has the secondary benefit of creating a record of who authorised expensive work, which is useful for cost attribution and for the investigation that follows an incident.

The internal version of this attack deserves mention because it is more common than the adversarial one and generates the same bill. A well-intentioned employee who asks an agent for a comprehensive analysis across everything the organisation holds has

issued an instruction with the same amplification properties as the attack described above, without any hostile intent. Deployments that treat expensive-plan confirmation as a security control aimed at attackers will find that it mostly catches colleagues, which is a good outcome and should be framed that way when the control is introduced, since the alternative framing invites objections about treating staff as suspects.

A final observation about the legal ambiguity noted above. The uncertainty cuts in both directions, and organisations should be aware of the second. An enterprise whose agent, having been manipulated, performs expensive operations against a third party's metered service may find itself the party generating the bill rather than receiving it. Agents that call external paid services on the organisation's behalf create an exposure that runs outward as well as inward, and the consumption ceilings described here should apply to outbound calls against third-party meters with the same discipline applied to the organisation's own model spend.

Rate of change is the alerting property that matters here, and it is worth restating in this context because the arbitrage attack has a distinctive profile. Ordinary agentic consumption is bursty but bounded: a task consumes, completes, and stops. An arbitrage attack produces sustained consumption at an elevated rate that does not terminate, because the plan the agent formed has no early exit. Alerting on the second derivative of spend, meaning acceleration rather than level, catches this within a minute or two of onset while a threshold on the absolute figure catches it after the budget is materially depleted.

A supplementary control that costs nothing is a schedule constraint. Most organisations have no legitimate reason for an agent to initiate a multi-million-operation task at three in the morning on a Sunday, and refusing to begin expensive plans outside defined windows removes the unattended overnight case that produces the largest individual incidents. The constraint is crude, it inconveniences the occasional legitimate batch job, and it converts the worst realistic outcome from an eight-hour burn into a bounded one.

13.3 Enterprise Alternatives & Structural Shifts

Every control described so far constrains consumption within a pricing model that makes consumption unbounded. A different approach changes the pricing model, and organisations have been moving in that direction during 2026 for reasons that are partly financial and, as this section argues, substantially architectural.

Deployment model	Mechanism	Effect on exposure
Hybrid credit pools	Flat subscription combined with a capped allowance of metered usage	Bounds the worst case at the cap; a runaway loop exhausts an allowance rather than generating an open invoice
Self-hosted open-weight models	Models operated on infrastructure the organisation owns or reserves	Decouples execution from variable cost entirely; a loop consumes capacity already paid for
Local small language models	Compact models running on edge or workstation hardware	Removes the external meter for the workloads they can serve, and removes the data egress with it

Table 13.1 Deployment structures and how each alters consumption exposure.

The security argument for these structures is stronger than the cost argument, and it is frequently overlooked because the discussion is usually led by finance. Under metered pricing, a runaway loop is a financial incident with no ceiling, developing silently and discovered on a bill. Under fixed capacity, the same loop is an availability incident: it consumes the capacity available, other work queues behind it, monitoring detects the saturation within seconds, and the response procedures are the ones every operations team has practised for a decade. The failure has been converted from an unfamiliar class into a familiar one, which is a substantial improvement in an organisation's ability to respond.

Data locality is a second benefit and is decisive in some regulated contexts. A locally hosted model processes content without transmitting it to an external service, which removes an entire category of question about where data went, who could access it, and under which jurisdiction it was processed. For organisations that have been unable to

deploy agentic capability against sensitive data for exactly these reasons, local models are the arrangement that makes deployment possible rather than merely cheaper.

The trade-offs are real and should be stated with the same clarity.

Capability is the first. Open-weight models have narrowed the gap against the leading proprietary systems considerably and have not closed it, and for tasks requiring the strongest available reasoning the difference is measurable. Small models running on local hardware are further behind again. The honest framing is that a portion of the workload can be served locally at no meaningful loss, a portion cannot, and determining the boundary requires evaluation against the organisation's actual tasks rather than against published benchmarks.

Operational burden is the second. Self-hosting means the organisation owns capacity planning, hardware procurement, model updates, and the reliability of the serving infrastructure. Teams that have run this comparison sometimes find the total cost of ownership exceeds the metered alternative at their volume, and that finding is legitimate. The security benefit described above may still justify the arrangement, and it should be argued explicitly rather than assumed.

Security responsibility is the third and is the one most often missed. Self-hosting moves the model serving infrastructure inside the organisation's boundary, which means the organisation now owns securing it. Model weights become an asset requiring protection. Inference endpoints become internal services requiring authentication and rate limiting. The shared responsibility diagram from Chapter 12 shifts, and everything that moves across the line becomes the customer's problem.

The arrangement most organisations arrive at is a tiered one, and it is worth describing because it is more defensible than any single choice. High-volume, well-defined, latency-sensitive tasks run on local or self-hosted models, where the workload is predictable and the capability requirement modest. Tasks requiring the strongest reasoning route to metered providers, with the consumption controls of Chapter 5 applied, because the volume is lower and the ceiling is therefore easier to enforce.

Sensitive data stays local regardless of capability requirements, which occasionally means accepting a worse result rather than a compliance exposure.

Routing between tiers becomes a security control in its own right under this arrangement, which is a point worth making explicitly because it is generally treated as a cost optimisation. The component that decides which model serves a request is deciding where data goes, what a task can cost, and which capability an attacker gains access to through a successful injection. It deserves the same review as any other enforcement point, and in several deployments it has been implemented as an afterthought in application code with no logging at all.

Evaluating where the boundary between tiers should fall is an exercise organisations tend to over-think, and a simple procedure produces most of the answer. Take a representative sample of completed agentic tasks, run each against the candidate local model, and compare outputs against what the metered provider produced. The proportion that are equivalent for the purpose at hand is the proportion of the workload that can move, and the comparison usually reveals that the distribution is bimodal: a large body of routine work where the difference is imperceptible, and a smaller body of genuinely difficult reasoning where it is decisive. Teams that reason about this in the abstract generally overestimate how much of their workload needs the strongest model available.

A caution about the direction of travel is warranted, since a book written at one moment can mislead a reader coming to it at another. Capability at every tier is improving, the gap between open-weight and proprietary systems has narrowed materially over successive generations, and hardware capable of serving useful models locally is becoming ordinary. A boundary drawn today will be wrong within a year, and it will be wrong in the direction of having moved too little work locally rather than too much. Building the routing layer so that the boundary is a configuration rather than an architecture is what makes revisiting it inexpensive.

Contract structure deserves a closing note in this chapter, since the deployment model and the commercial arrangement are separable and organisations frequently conflate them. It is possible to consume a proprietary model under a committed-capacity agreement rather than a pure meter, which supplies a fixed cost base without operating any infrastructure. It is equally possible to self-host and still face variable cost, where the hosting is itself metered. The property that matters for the exposure described in this chapter is whether the marginal cost of an additional model call is zero or positive, and that is a commercial question rather than an architectural one. An organisation that wants the security benefit of fixed cost should negotiate for it directly before concluding that self-hosting is the only route to it.

Reserved capacity carries its own failure mode that should be understood rather than discovered. Under a fixed capacity arrangement, a runaway execution consumes shared capacity, and the workloads competing for it degrade. This is the availability incident described above, which is a considerable improvement on an unbounded bill, and it is still an incident. Per-workload capacity reservations, so that one runaway cannot starve the rest, are the equivalent of the per-task budget ceiling in the metered world, and they should be configured for the same reason.

The next chapter converts the arguments of this book into a sequenced plan, and this chapter's contribution to it is a specific one: consumption controls belong in the first tier of work, not because the financial exposure is the largest an organisation faces, but because they are among the cheapest controls to implement and the only ones that bound a failure mode which needs no adversary to occur.

Part Six: Architectural Blueprint for Executive Action

Chapter 14: Architectural Blueprint for Executive Action and Long-Term Strategy

A book of this kind owes its reader a sequence rather than a list. The preceding chapters have described a large number of controls, and an organisation attempting all of them simultaneously will complete none. This chapter sets out what to do first, why that order and not another, and which of these problems will still be problems when the current generation of tooling has matured.

14.1 Immediate Technical Directives for CISOs

Five directives follow, ordered by dependency rather than by importance. Each is a prerequisite for those beneath it in the sense that attempting a later item without the earlier ones produces a control with insufficient input to operate on.

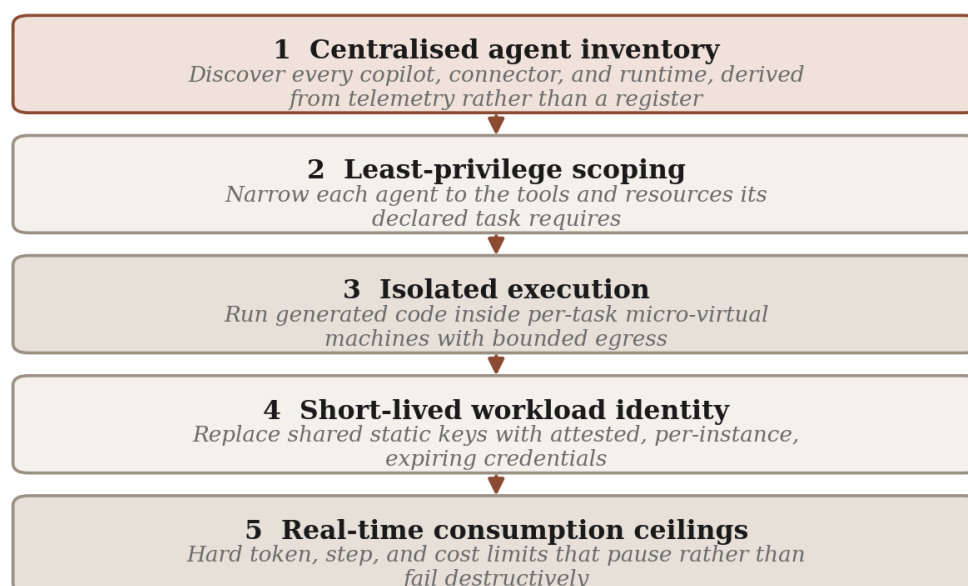


Figure 14.1 Five directives in dependency order. Each earlier step supplies input the later steps require.

Build a centralised agent inventory. This is first because nothing else is possible without it, and because most organisations do not have one. The inventory must record every agent, meaning every orchestrator deployment, every activated vendor feature, every

developer tool that executes commands, and every browser-resident assistant acting inside authenticated sessions. For each, it should record the tools available to it, the credentials it presents, the data it can reach, the principal it acts for, and the team accountable for it.

Two properties separate an inventory that works from one that decays. It must be derived from runtime telemetry rather than maintained by hand, because agents are created by software and a manually maintained register is wrong within days. And it must extend to the doors described in Chapter 1 that are not the platform deployment: the vendor feature enabled by default, the coding assistant on a developer's machine, the browser extension operating with a user's full entitlements. An inventory covering only the orchestration platform documents the population that was already visible.

Enforce least-privilege scoping. With an inventory in place, the second directive narrows what each agent can reach. The practical method is the one described in Chapters 8 and 10: derive an initial scope from observed behaviour, have the owning team review and adopt it, operate default-deny on anything undeclared. Scope by resource rather than by operation where possible, since restricting an agent to the records its task concerns is a stronger constraint than restricting it to read operations across everything.

The classification that makes this tractable is reversibility, which has appeared throughout this book and should be built once and centrally. Every tool in the registry is marked according to whether its effects can be undone. Reversible actions are permitted with logging. Irreversible actions require narrow scope, confirmation, or both. This single classification drives approval gates, failure modes for the enforcement layers, alert priority, and the tiering of every semantic control, and an organisation that maintains it once has supplied the input for all of them.

Isolate execution environments. Where agents execute generated code, that execution belongs inside per-task micro-virtual machines with ephemeral storage, restricted system call surfaces, and outbound access limited to an allowlist. The instance is destroyed on task completion. Agents that call only defined tool interfaces and execute

no code derive little from this and should not be made to pay its cost, which makes the distinction between the two populations part of the work.

Verification matters as much as configuration here. Attempting, from inside the sandbox, to reach an unlisted endpoint, to write outside the scratch space, and to use the instance credential outside the task scope takes an afternoon and regularly reveals that one of the constraints is not actually enforced.

Transition to short-lived workload identities. Shared static keys are the credential pattern with the worst properties available, and roughly half of surveyed organisations still use them for agent-to-agent communication. Replacing them with attested, per-instance, expiring credentials restores attribution, permits scope narrowing, makes revocation automatic, and bounds lateral movement.

The sequencing within this directive determines whether it delivers value early or only at completion. Distinct identities per agent come first and restore attribution immediately, even where the credentials remain long-lived. Narrow scopes follow. Short lifetimes follow those. Attestation completes the architecture. An organisation that designs the full model, discovers it requires infrastructure it does not have, and ships nothing has done worse than one that issued distinct credentials in a fortnight.

Implement real-time consumption ceilings. Budget limits per task, step limits per session, and semantic drift termination, enforced by a component positioned between the orchestrator and the model provider. Ceilings should pause and require authorisation rather than failing destructively, since a control that halts legitimate work without recourse will be disabled by the first team it inconveniences. Alerting should trigger on rate of change rather than on absolute level, because agentic consumption moves faster than a threshold crossing can be delivered.

This directive is last in dependency order and should not be last in execution order. It is the cheapest of the five, it requires no organisational change, and it bounds a failure mode that occurs without any adversary. Organisations with limited capacity should

implement inventory and consumption ceilings in parallel and take the middle three in sequence.

Funding the programme is a question every security leader reading this will face, and the framing that has worked is worth passing on. Presented as agentic security work, this programme competes for budget against other security initiatives with clearer incident histories. Presented accurately, most of it is platform work with security benefits: the inventory is an asset management problem, the identity work delivers cost attribution and debugging improvement, the consumption ceilings are a financial operations control, and the telemetry serves reliability engineering as much as investigation. Several organisations have funded the whole of it from platform and financial operations budgets, and there is no reason for a security architect to be precious about which cost centre pays for controls they need.

Sequencing against business appetite matters as much as sequencing against dependencies, and this is where technically sound plans fail. A programme that arrives at a business unit demanding scope declarations before that unit has any working agentic capability will be resented and evaded. A programme that arrives offering a paved path, meaning an approved deployment route with credentials, logging, and connectors already configured, is adopted because it is faster than building privately. The controls travel with the path. This is the same argument Chapter 1 made about shadow deployment, and it generalises: the security function gains ground here by making the compliant route the convenient one rather than from prohibition it cannot enforce.

Two measurement questions accompany the plan and are worth settling early, because a programme that cannot describe its own progress loses funding regardless of its technical merit. The first is coverage: what proportion of the known agent population is subject to each control. This is straightforward to compute once the inventory exists and is the single most useful number a security leader can carry into a board discussion, because it converts an abstract concern into a fraction that can be moved. The second is

exposure concentration: what proportion of agents hold irreversible capability, and what proportion of those are fully covered. A programme reporting ninety percent coverage while the uncovered ten percent contains every agent that can move money has reported a number that conceals its own position.

Neither measurement requires new tooling, and both are derivable from the inventory and the reversibility classification. What they require is the discipline to compute them monthly and to report the second alongside the first, since the first alone is the number that makes a programme look healthy and the second is the one that tells its leader where the remaining risk sits.

A realistic timeline helps set expectations, and the figures that follow reflect what organisations of moderate size have reported rather than any theoretical estimate. Inventory takes six to twelve weeks, most of it spent on discovery of the deployments nobody registered. Consumption ceilings take two to four weeks and can proceed in parallel. Scoping takes a quarter or more because it depends on business teams producing declarations. Isolation takes a quarter where the platform already runs container workloads and considerably longer where it does not. Identity is the longest item, typically two to three quarters, because it touches every resource server. An organisation planning on a twelve-month horizon can complete all five; one planning on a quarter should do inventory and ceilings and be honest that the rest is scheduled rather than done.

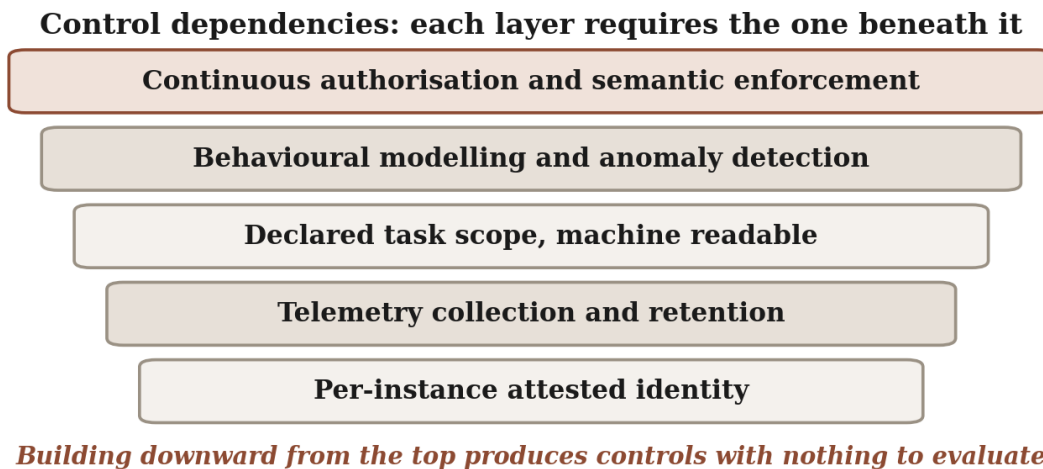


Figure 14.2 Control dependencies. Each layer supplies the input the layer above it evaluates.

A note on what to do about the agents already running while this programme is under way, since the plan above reads as though the estate can be paused. It cannot, and the interim posture matters. The controls that can be applied to an existing deployment without redesigning it are worth applying immediately: distinct credentials in place of shared ones, consumption ceilings at the provider account, logging of tool calls where the framework supports it, and removal of tool grants that nobody can justify. None of these requires the full architecture, each takes days rather than quarters, and together they address a meaningful share of the exposure while the durable work proceeds.

Prioritising among existing agents is the other half of the interim question, and the reversibility classification answers it. Agents holding irreversible capability, meaning those that can move money, publish externally, delete records, or alter permissions, should receive attention first and disproportionately. In most estates this is a small minority of the population, frequently under a tenth, which makes the first pass tractable in a way that securing everything is not. An organisation that has properly constrained the agents that can cause unrecoverable harm has bought itself the time to do the rest properly.

A note on skills, since the plan assumes capabilities that most security functions do not currently hold. Reading an execution trace and forming a judgement about whether an

agent's reasoning was manipulated is unlike any existing security discipline; it resembles fraud analysis more than intrusion analysis. Writing scope policy requires understanding both the business process and the tool interface. Evaluating a model routing decision requires knowing what the models can do. None of these is exotic, and all of them are more easily developed inside the team than hired for, because the external market barely exists. Organisations that have made progress generally seconded an engineer into the security function rather than attempting to recruit a specialist.

Relationships with the platform team determine whether any of this works, and the dependency runs one way. Every control in this book requires something from the orchestration layer: scope declarations, telemetry emission, delegation records, identity integration, and a hook that tells the detection layer when an agent's configuration changed. None of it can be built by the security function alone, and a security programme that specifies requirements at the platform team rather than with them will find the requirements deprioritised. The organisations that have moved fastest embedded a security engineer in the platform team and let the requirements arrive as design input rather than as an audit finding.

Three organisational changes accompany the technical work, and they are harder to make than any of the five directives.

Ownership has to be assigned explicitly. Agents are deployed by business units, run on platforms owned by engineering, call models operated by vendors, and touch data owned by other functions. During an incident every one of those parties has to act, and none reports to the others. Establishing in advance who may stop an agent, who owns its memory store, and who is accountable for its scope is a governance exercise rather than a technical one, and organisations that have handled incidents well did this before the incident.

Change control has to extend to prompts and tool registries. A system prompt determines what an agent does, which makes it functionally equivalent to source code,

and in most deployments it is edited through a console without review, versioning, or an audit record. Bringing it under existing change management requires no new control, only the recognition that it is a configuration item.

Security has to be present where agents are procured, which is a different set of conversations from the ones security is usually invited to. The decisions that determine exposure are made when a vendor feature is enabled, when a service agreement is signed without observability commitments, and when a team selects an orchestration framework. By the time an architecture review occurs, the constraints are set.

Two things this book has argued for do not appear in the five directives, and their absence is deliberate. Semantic firewalls and continuous authorisation are the most sophisticated controls described here and the most demanding to operate. They require identity, telemetry, scope declaration, and a policy corpus that someone maintains. An organisation that has completed the five directives is positioned to add them and will find the addition straightforward. An organisation that begins with them, which is tempting because the vendor offerings are clearest at that layer, will have purchased a component with nothing to feed it.

A word on where to place effort when the estate contains vendor-managed agentic services that the organisation cannot instrument. Several of the directives above assume control over the runtime, and where capability is consumed as a service that assumption fails. What remains available is the boundary: which data the service is given, which of the organisation's systems it may reach, which credentials it presents, and what it is permitted to do with the results. Constraining a service at its edges is weaker than constraining an agent internally, and it remains the whole of what an organisation controls in that arrangement. The corresponding action is to make the observability and containment questions of Chapter 12 part of procurement, so that the next such service arrives with better properties than the last.

Reporting to a board deserves a paragraph because the framing determines whether the programme survives a budget cycle. Boards respond poorly to catalogues of technical

controls and well to a small number of questions with tracked answers. The four questions this chapter closes with serve that purpose directly: which agents exist, what each may do, what each has done, and who can stop them. Reporting the proportion of the estate for which each question can be answered, with the proportion for irreversible-capability agents shown separately, gives a board something it can govern. It also has the property of being honest, which control-count reporting generally is not.

Common failure patterns in these programmes are worth naming so that a reader can recognise their own, and four recur often enough to be predictable. The first is starting at the top of the dependency stack, usually because the vendor offerings are clearest there, and ending with an enforcement component that has no scope declarations to enforce and no telemetry to evaluate. The second is securing the visible deployment while the vendor-enabled features, developer tooling, and browser-resident assistants remain outside the inventory entirely, which produces a coverage figure that is accurate about the wrong population.

The third is treating the programme as a project with a completion date. Agent populations grow, configurations change weekly, and a control set validated in March describes a system that no longer exists in September. Continuous posture evaluation, meaning ongoing comparison of the current configuration against the intended one, is the mechanism that keeps a completed programme from decaying, and it is the item most often cut when the project is declared finished. The fourth is building the technical controls without settling ownership, which produces an organisation that can detect an agentic incident and cannot decide who is permitted to stop it.

Each of these has the same root, which is that the work is easier to specify technically than to sustain organisationally. The five directives can be described in a page. Keeping them true across a growing estate, a changing platform, and a business under delivery pressure is a different activity, and it is the one that determines whether the programme is worth anything two years after it was funded.

Looking further out, three of the problems described in this book are engineering problems and will be solved. Identity for autonomous workloads is a solved architecture awaiting integration. Telemetry standardisation is under way and will settle. Consumption control is straightforward once someone builds it, and the components exist.

Two problems are structural and will not be engineered away. The first is that instructions and data occupy the same channel in a language model, which is a property of how these systems work rather than a defect in how they are built. Mitigations will improve and the boundary will remain probabilistic, which means architectures must continue to assume that an agent's objective can be redirected and place their controls at the point of action accordingly. The second is that autonomy and control are in tension by definition. Every constraint described in this book reduces what an agent can do, and the value of the technology comes from what it can do without being asked. Organisations will keep negotiating that trade, and the negotiation will keep being settled under commercial pressure rather than on the security team's terms.

A more useful posture, given both, is to stop treating agentic security as a problem with a completion state. The question is not whether an organisation has secured its agents. It is whether it knows which agents exist, what each may do, what each has done, and who can stop them. An organisation that can answer those four questions is in a position to make sound decisions about everything else, and one that cannot is making decisions in the dark regardless of how many controls appear on its inventory.

Conclusion

This book began with a claim about obedience. Enterprise software was for most of its history obedient in a narrow sense: it did what its authors compiled it to do, and when it failed it failed along paths an engineer could trace to a line of source. The security architecture the industry built assumed that property everywhere, in firewalls that filtered predictable traffic, in scanners that inspected the code that would run, and in access control that granted permissions to identities attached, eventually, to a person who could be telephoned.

Fourteen chapters later the claim can be stated more precisely. What autonomy removed was not the obedience itself, since agents are if anything more obedient than the software they replaced. What it removed was the guarantee about whose instructions they obey. An agent follows the instructions in its context, and its context contains text from sources the organisation does not control. The compliance that makes these systems useful is the same property that makes them exploitable, and no amount of tuning separates the two because they are the same mechanism observed from different sides.

Three claims were made in the Introduction, and each has been developed at length. It is worth restating them in the form the intervening material has given them.

Security boundaries in these systems are semantic rather than syntactic. Every consequential attack described in this book proceeds through requests that are valid in every mechanical respect: correct protocol, valid credential, permitted destination, authorised operation. Chapter 1 explained why the three main categories of enterprise defensive tooling are structurally incapable of seeing this, Chapter 7 explained why behavioural baselines built for humans fail against continuous machine workloads, and Chapters 8 and 10 described the enforcement points that can hold the necessary context. The recurring finding across all of them is that the question a control must answer has changed from whether a request is well formed to whether it serves the task the agent was given, and that the second question requires state the existing stack does not carry.

Autonomy converts theoretical vulnerabilities into practical ones. The confused deputy was a curiosity for thirty years because exploiting it required an unusual arrangement of privileges. Chapter 2 showed that an agent with tool access is a confused deputy by construction, sitting permanently in production, taking direction from whatever text it retrieves. Chapter 6 showed the same reasoning applied to internal communication, where trust assumptions borrowed from service meshes turn out to be badly wrong. The NIST measurement cited in Chapter 2, an eighty-one percent attack success rate against agentic systems compared with eleven percent against conventional controls, is roughly what the structural argument predicts, and the reason the ratio is so wide is that the attacker's task has changed from defeating a control to persuading a component.

Economic exposure has widened alongside the technical one. Chapters 5 and 13 treated consumption as a security concern rather than a financial one, because an attacker who can trigger recursive execution imposes cost without breaking anything, and because the resilience engineers deliberately design into these systems is the mechanism that converts a single sentence into an unbounded invoice. This is the one exposure in the book that requires no adversary at all, which is why the consumption controls appear early in the plan despite being the least sophisticated.

Something should also be said about what this book has asked of the reader's organisation, since the aggregate is substantial and a book can propose more than any team can absorb. The five directives of Chapter 14 are the compressed form. Beneath them sit the arguments that produced them, and the value of having read the arguments rather than only the list is that the list will need adapting. An organisation whose agents execute no generated code can skip most of the isolation work. One whose agentic estate is entirely vendor-managed has a procurement problem rather than an engineering one. One operating under sectoral supervision will find the compliance obligations set its retention windows rather than the detection interval. The reasoning is what makes those adaptations sound; the list alone would produce a programme that fits nobody's circumstances precisely.

Two developments in the attacker's position deserve restating alongside those claims, because Part II described a shift that is easy to lose among the defensive material. The first is that the trade-off between scale and precision, which shaped computer-enabled fraud and social engineering for three decades, has dissolved. Research that once justified only high-value targets now costs almost nothing, which means the population worth attacking individually has expanded to include anyone whose access has any value. Awareness programmes built around detecting the artefacts of cheap attacks are teaching people to look for signals that no longer appear.

The second is that the attacker's activity has become indistinguishable from ordinary use at every layer that examines individual events. A synthetic identity matured over weeks behaves like a customer. A runaway execution behaves like a difficult task. A cascading payload behaves like an internal message. Each of these defeats point-in-time evaluation for the same reason, which is that the property distinguishing them from legitimate activity exists in a relationship across events rather than in any event considered alone. That is the same finding as the semantic boundary claim, arriving from the attacker's side, and its practical consequence is identical: detection has to model sequences and relationships or it will model nothing useful at all.

Underneath the specific findings, one requirement surfaced from four independent directions, and its recurrence is the most useful thing this book has to offer. Identity work in Chapter 3 required knowing what a session was authorised to do. Behavioural modelling in Chapter 7 required a notion of intended behaviour to compare observed behaviour against. Semantic enforcement in Chapter 8 required a scope to evaluate requests against. Continuous authorisation in Chapter 10 required a declared task to evaluate actions within. The standards work in Chapter 11 arrived at the same place independently and gave it a name, the machine-readable capability manifest. Every one of these is the same requirement: an agent's permitted behaviour has to be written down in a form a machine can evaluate. An organisation that has not done that cannot implement any of the controls in this book, and one that has done it will find several of them straightforward.

A second thread runs through the book alongside that one and deserves the same treatment. Repeatedly, the analysis has arrived at a distinction between actions that can be undone and actions that cannot, and has used it to allocate scarce protective effort. Approval gates belong on irreversible actions. Enforcement controls fail closed for irreversible actions and open for the rest. Alert priority follows it. Tiering of expensive semantic evaluation follows it. Interim prioritisation across an unsecured estate follows it. Five separate controls in this book depend on the same classification, which means an organisation should build it once, centrally, as a property of the tool registry rather than as a judgement each control makes independently.

That classification is also the most defensible thing a security function can offer a business under delivery pressure. A blanket demand to constrain every agent invites a negotiation the security function will lose. A demand to constrain the small minority of agents that can move money, delete records, publish externally, or alter permissions is proportionate, comprehensible to a non-technical executive, and generally accepted. In most estates that minority is under a tenth of the population, which makes the ask tractable and the refusal difficult to justify.

Convergence on that requirement supplies the practical answer to the question a security leader is likely to arrive at this page holding, which is where to start when everything appears to be a prerequisite for everything else. Start with the inventory, because nothing can be scoped before it is known to exist. Then declare, for each agent that matters, what it is permitted to do. Everything else in this book attaches to those two artefacts.

It is worth recording which of the arguments here rest on firm ground and which on inference, since a reader will apply them under conditions this book cannot anticipate. The structural claims are firm: the absence of an instruction and data boundary in language models, the confused deputy property of tool-calling agents, the failure of user-centric baselining against continuous workloads, and the coupling of correctness to

cost under metered pricing all follow from properties of the architecture and would hold even if every statistic in this book were withdrawn.

The quantitative claims are weaker and were presented as such. Survey figures on deployment, governance, and credential practice come from vendor-commissioned research with self-selected respondents and inconsistent definitions. Incident cost estimates are assembled from spending patterns rather than measured. Red-team success rates depend heavily on how success was defined. Each of these was used to establish an order of magnitude and a direction, and the arguments were constructed to survive substantial error in any of them. A reader who finds that a specific figure has been revised should check whether the argument depended on the number or only on its sign; in nearly every case it is the latter.

The forward-looking claims are weakest of all, and the ones about how standards and regulation will settle should be treated as informed guesses. What can be said with more confidence is that the requirements which recur across every framework, meaning attribution, reconstructibility, and demonstrable human authority over consequential actions, are recurring because they follow from the architecture rather than from any drafting choice, and they are therefore the safest thing to build toward regardless of which specific instrument prevails.

A word about what this book has not settled. The prompt injection problem has no solution and will not acquire one, because the property that produces it is the property that makes language models useful. The mitigations described in Chapter 2 raise the cost of an attack and do not create a boundary, and any architecture that treats them as a boundary has misunderstood what it bought. The correct response is the one this book has argued for throughout, which is to move the load to the point of action: if a redirected agent cannot perform a damaging operation, the injection succeeds semantically and fails operationally, and that is the only outcome available with a guarantee attached.

Tension between autonomy and control is likewise unresolved and structural. Every constraint in these chapters reduces what an agent can do, and the value of the technology lies in what it does without being asked. Organisations will keep negotiating that trade, and the negotiations will keep being decided under commercial pressure. A security function that treats each such decision as a failure will exhaust itself. The more durable posture is to make the trade explicit, record what was accepted and by whom, and concentrate the non-negotiable constraints on the small set of actions that cannot be undone.

Three developments were set out in the Introduction as things that would falsify the argument of this book, and they are worth leaving with the reader as things to watch rather than as rhetorical devices. A durable architectural separation between instructions and data at the model layer would make the injection chapters historical. A demonstration that semantic evaluation can be performed cheaply and reliably at production volume would make much of the session-scoped machinery here unnecessary. Sustained growth in deployment without a corresponding growth in incidents would suggest that something about real systems constrains these attacks in a way the structural analysis missed. Any of the three would justify reallocating investment, and a reader encountering this material some years after it was written should weigh the incident record more heavily than the reasoning.

None of what precedes is an argument against deploying autonomous systems, and it would be a poor reading of this book to take it as one. The productivity case is real, deployment is proceeding at a rate no security function will reverse by memorandum, and a programme built on the premise that the technology should not exist will be routed around within a quarter. The question that has organised these chapters is narrower and more useful: what has to be true about an agentic deployment before a competent professional can defend it in front of a board.

A final substantive point about the shape of the field, offered because it affects how a reader should treat everything in these pages. Agentic security is currently a discipline

in which practice runs ahead of published evidence. The mechanisms are well understood, the incident record is thin and biased toward what was visible from outside, and the control patterns described here reflect what practitioners have found workable rather than what controlled study has demonstrated effective. That is an uncomfortable basis for architecture and it is the basis available, and stating it plainly is better than the alternative convention of presenting inference as finding.

What follows practically is a preference for controls whose value does not depend on the threat model being correct. Attribution is worth having whether or not agentic attacks turn out to be common, because attribution serves debugging, cost allocation, and compliance. Telemetry is worth having for the same reasons. Scope declaration improves reliability independent of any security benefit, since an agent constrained to what it needs fails more predictably. Consumption ceilings prevent accidental overspend whether or not anyone attacks. Each of the five directives has a justification that survives the security argument being overstated, which is the property to look for when building under uncertainty.

A note on the people who will do this work. Nothing in these chapters is exotic engineering. Workload attestation, token exchange, policy evaluation, structured telemetry, and consumption metering are all mature technologies, and the difficulty in an agentic programme almost never comes from a missing component. It comes from the components sitting in different teams, that the integration between them is nobody's assigned responsibility, and that the deadline for the agentic capability was last quarter. Security functions that have made progress did so by treating the problem as one of coordination rather than of technology, which is unglamorous and is where the actual obstacle sits.

That framing also suggests where a small team should concentrate. Given limited capacity, the highest-value activity is establishing a fact rather than building a control: which agents exist, what each can reach, and who owns it. Facts of that kind are cheap to gather while an estate is small, expensive once it is large, and prerequisite to everything

else. A security function that has spent a quarter establishing them and building nothing has done better than one that deployed a sophisticated control across the fraction of the estate it happened to know about.

There is a broader observation available here about how security disciplines mature, and it offers some grounds for confidence. Every previous shift in enterprise computing produced a period in which deployment outran control and the tooling lagged the architecture. Virtualisation, cloud, containers, and continuous deployment each arrived faster than the security practice around them, each generated a body of incidents that looked unprecedented at the time, and each eventually acquired a control model that practitioners now consider ordinary. The current moment resembles the early period of each of those, with the important difference that the pace is faster and the failure modes involve a component whose behaviour cannot be fully specified.

What follows from that comparison is a specific piece of advice rather than reassurance. In each previous case, the organisations that came through well were the ones that built the primitives early: identity, inventory, and telemetry, before the sophisticated controls that depend on them, and while the estate was still small enough to instrument. The organisations that struggled were the ones that waited for the tooling to mature and then attempted to retrofit attribution onto a population that had grown past the point where it could be enumerated. That pattern is repeating now, the window in which an agent estate is small enough to inventory cheaply is closing, and the argument for acting during it is the same one that held in every previous cycle.

The answer, compressed to its shortest defensible form, is four questions. Does the organisation know which agents exist. Does it know what each is permitted to do. Does it know what each has actually done. Does it know who can stop them. An organisation that can answer all four is equipped to reason about everything else in this book and to make sound judgements about the trades it will inevitably have to accept. An organisation that cannot answer them has no position from which to make decisions about agentic security at all, whatever appears on its control inventory. Those four

questions are where the work starts, and the rest of it follows from having answered them honestly.

Acknowledgments

A book about systems that act without supervision is, fittingly, the product of a great deal of supervision. Whatever is useful here came out of arguments with people who knew more than I did about some part of it, and whatever is wrong survived those arguments through my own stubbornness.

My thanks go first to the engineers and architects who talked me through deployments they had every commercial reason to keep quiet about. Several of the failure patterns described in these chapters exist in the public record only because somebody was willing to describe an uncomfortable morning in detail. That kind of candour is how a young field learns anything at all, and it is rarer than it should be.

I am grateful to colleagues who read early drafts and told me plainly where the argument did not hold. Chapters 7 and 10 in particular are considerably better for objections I did not want to hear, and the sections acknowledging what this book cannot settle exist because people insisted that a confident answer would have been a dishonest one.

The researchers, standards contributors, and security teams whose published work is cited throughout deserve particular acknowledgment. A field this new depends on the people prepared to measure something and publish the result, including when the result is inconvenient. Where I have disagreed with their conclusions, I have tried to do so on the evidence and to make the disagreement visible.

My family absorbed the cost of the evenings and weekends this took, without complaint they were entitled to make. That patience is the reason the book exists in a finished form rather than as a folder of notes.

Any errors of fact, judgement, or emphasis that remain are mine alone.

About the Author

Khushan Adatiya is a Senior Software Engineer working on large-scale distributed systems, with more than thirteen years of experience building platforms that operate at global scale. His technical focus spans distributed systems architecture, artificial intelligence, big data and data engineering, and cloud observability, security, and governance.

Across his career he has worked on infrastructure sustaining hundreds of thousands of transactions per second, concentrating on the observability, security, and governance layers that determine whether systems of that size remain understandable to the people responsible for them. His current work centres on artificial intelligence and search infrastructure, following earlier roles in cloud artificial intelligence platforms and large-scale cloud services.

He is a Senior Member of the IEEE and a member of Raptors.dev and SigmaXi. He serves regularly as a peer reviewer for academic venues, including ICMI Late Breaking Results and the ICMI Grand Challenge, UbiComp and ISWC, Insights, and the Journal of Open Research Software. He has also judged a range of professional technology competitions and awards programmes alongside hackathons including the Gemini Live Agent Challenge, the Hack Nation Global AI Hackathon, the Global Fusion Hackathon, and Developer Week.

He writes regularly on distributed systems, agentic architecture, and the operational economics of artificial intelligence. His articles have appeared in HackerNoon, covering agent-to-agent protocols, agents delivered as a service, production scalability, and the shift toward agent-native systems, and in AI Journal. His work on observability and its cost implications for the API economy has been covered by TechBullion, and he has been profiled in an interview with IndieHackers.

Khushan holds a Bachelor of Engineering in Computer Engineering and a Master of Science in Computer Science. He lives in Fremont, California.

Securing the Agentic Enterprise is his first book.

Glossary

Terms are defined as they are used in this book. Where usage in the wider literature is unsettled, the definition here is the narrower one, and the chapter in which the term is developed is noted.

A2A (agent-to-agent) communication. Messages exchanged directly between autonomous agents, typically carrying a delegated subtask, its context, and a claim of authority. Treated in this book as untrusted input regardless of origin. See Chapter 6.

Agent. A software process that receives an objective in natural language, decomposes it into steps without external direction, selects and invokes tools, and evaluates results before deciding what to do next. The three load-bearing properties are an under-specified objective, a plan produced at runtime, and a loop that continues without per-step approval.

Agent passport. A short-lived credential that cryptographically binds an agent instance, the task it is performing, and the human principal whose authority it acts under, so that a downstream resource can verify the whole delegation path. See Chapter 3.

Agent sprawl. Growth in the deployed agent population beyond what an organisation has inventoried or reviewed, typically because business units provision agents directly. See Chapter 1.

Attenuating delegation. A property of a delegation chain in which each hop may narrow authority and never widen it, so that an agent several hops from the origin holds strictly less authority than the originating principal.

Confused deputy. A program holding privileges of its own that can be induced by a less privileged caller to exercise those privileges on the caller behalf. Described by Norm Hardy in 1988; an agent with tool access is one by construction. See Chapter 2.

Continuous authorization. Evaluation of authority before every action rather than once at session start, so that a request decided after the token was issued can still be refused. See Chapter 10.

Cost arbitrage attack. Supplying an instruction whose execution is legitimate but expensive, causing the target own authorised systems to consume its budget. See Chapter 13.

Dark agent traffic. Inter-agent communication that exchanges context, delegations, and credentials outside any collection point the security function can query. See Chapter 8.

Denial of wallet (DoW). An attack whose objective is financial exhaustion rather than resource exhaustion, forcing a target metered systems into high-cost processing paths. See Chapter 5.

Ephemeral execution environment. A container or micro-virtual machine provisioned in response to an agent plan and destroyed on task completion, so that anything written inside it does not persist.

Execution graph. The runtime structure of an agentic task: the sequence of states an agent moved through and the tools it invoked at each. The object that transition-state detection models.

Indirect prompt injection. Placement of instructions in content an agent will retrieve in the ordinary course of its work, so that the model processes them as directives. The attacker never touches the target system. See Chapter 2.

Irreversibility. Whether an action can be undone. Used throughout this book as the primary criterion for allocating protective effort, ahead of the read and write distinction that access control systems expose.

Machine-readable capability manifest. A structured declaration of what an agent is permitted to do, evaluable by a system rather than described in documentation. Required independently by identity, detection, enforcement, authorisation, and standards work.

MCP (Model Context Protocol). A common connector standard between agents and the tools they call. Addresses connection and message format; carries no native mechanism for constraining an authenticated caller. See Chapter 6.

Memory injection. Insertion of a durable-sounding instruction into an agent persistent context store, where it remains latent until a matching condition retrieves it. See Chapter 2.

Micro-virtual machine. A lightweight virtualised runtime with its own kernel and a small host interface, startable in tens of milliseconds. Used for per-task isolation of agent code execution.

Non-human identity (NHI). A credential-bearing identity belonging to software rather than a person. Outnumber human identities in enterprise environments by roughly eighty to one.

Orchestrator. Software that maintains agent state, holds the tool registry, manages the execution loop, and coordinates multiple agents. The only component positioned to emit reasoning telemetry, and therefore treated in this book as security infrastructure.

Persistent memory. State deliberately retained across sessions, typically in a vector store or structured summary. A separate asset class from session context because its security properties differ in kind.

Prompt trace. The complete lineage of an agentic task: objective received, plan produced, reasoning at each step, tool observations, and plan revisions. The record that explains why an action was taken. See Chapter 9.

Provenance (of a memory entry). A record of which session wrote a retained entry, from which source, under which task, and when. The difference between a remediable incident and rebuilding a memory store.

Resilience trap. The pattern in which recovery behaviour that is individually sensible compounds into an unbounded and increasingly expensive sequence. See Chapter 5.

Semantic drift analysis. Measuring similarity between consecutive execution steps to distinguish progress from repetition. Used both to terminate runaway loops and to locate the step at which an injection took effect.

Semantic firewall. A mediating component in the inter-agent path that resolves the requested action from a message and evaluates it against the declared scope of the session before forwarding. See Chapter 8.

Semantic privilege escalation. Redirection of authority that already exists, rather than acquisition of new authority. A permissions audit afterwards finds nothing out of place. See Chapter 2.

Session. A bounded execution context beginning when an agent receives an objective and ending when it is completed or abandoned. The natural unit at which scope is declared, authority delegated, and behaviour baselined.

Session smuggling. Manipulation of a trusted orchestrator task graph by a compromised or malicious sub-agent, propagating influence upward through a hierarchy rather than downward.

Shadow AI. Agentic capability provisioned outside formal security or technology review. Distinguished from earlier unsanctioned software by being continuously active rather than idle between uses.

SPIFFE and SPIRE. A standard and reference implementation for issuing short-lived cryptographic identity documents to workloads on the basis of attested runtime properties rather than a stored secret.

Task scope. The machine-readable statement of what a session is permitted to do, declared at session start. The artefact on which identity narrowing, behavioural comparison, semantic enforcement, and per-action authorisation all depend.

Token exchange. An OAuth 2.0 mechanism by which a party holding one token obtains a second, narrower token representing a specific delegated authority. The basis of attenuating delegation between agents.

Transition-state detection. Modelling an agent behaviour as a sequence of states and treating an unmodelled transition between two individually authorised operations as the anomaly. See Chapter 7.

Workload identity. An identity belonging to a running instance, attested from its properties at runtime, rather than a secret held in configuration.

Works Cited

Sources consulted in the preparation of this book. Web addresses were current at the time of writing; readers should expect some to have moved.

1. Agentic AI Security: A CISO Playbook for 2026 - Simbian AI. <https://simbian.ai/blog/agentic-ai-security>
2. The Complete Anatomy of AgentOS: How the AI Agent Operating System Is Rewriting the Rules of Enterprise - note. <https://note.com/betaitohuman/n/nf36b85483d60>
3. The Complete Guide to AI Agent Security for Enterprises (2026) | NeuralTrust. <https://neuraltrust.ai/blog/ai-agent-security-enterprises-complete-guide>
4. Agentic AI Security Glossary. <https://salt.security/glossary>
5. What Is Workstation Agent Security? The 2026 Enterprise Defense Stack - Repello AI. <https://repello.ai/blog/workstation-agent-security>
6. What a Hardened MCP Server Looks Like - Salt Security. <https://salt.security/blog/the-mcp-security-blueprint-what-a-hardened-mcp-server-looks-like>
7. Agentic AI Risks: A Guide to Proper AI Governance - Strata.io. <https://www.strata.io/blog/agentic-identity/agentic-ai-governance-how-to-approach-it/>
8. Everyone is Deploying AI Agents. Who is Governing Them? - Gspann.com. <https://www.gspann.com/insights/blog/everyone-is-deploying-ai-agents-who-is-governing-them>
9. AI Agent Security in 2026: What Enterprises Are Getting Wrong - AGAT Software. <https://agatsoftware.com/ai-agent-security-enterprise-2026/>
10. AI Agent Security: Meta's Rogue AI Agent Incident - AGAT Software. <https://agatsoftware.com/ai-agent-security-meta-rogue-agent-incident/>
11. Shadow AI explained: risks, costs, and enterprise governance - Vectra AI. <https://www.vectra.ai/topics/shadow-ai>
12. Shadow AI Agents: How Autonomous Workflows Are Bypassing Your Data Security Controls in 2026 - Kitecyber. <https://www.kitecyber.com/shadow-ai-agents-data-security/>
13. Securing AI agents: the defining cybersecurity challenge of 2026. <https://www.bvp.com/atlas/securing-ai-agents-the-defining-cybersecurity-challenge-of-2026>
14. AI Agent-to-Agent Communication: The Next Major Attack Surface - Salt Security. <https://salt.security/blog/ai-agent-to-agent-communication-the-next-major-attack-surface>
15. Top Agentic AI Security Threats in Late 2026 - Stellar Cyber. <https://stellarcyber.ai/learn/agentic-ai-security-threats/>
16. Why IDC Believes the Next Billion Users Aren't Human (And How to Secure Them) - Blog. <https://www.menlosecurity.com/blog/why-idc-believes-the-next-billion-users-arent-human-and-how-to-secure-them>
17. From Browsers as a Risk to Browsers as Security Control Planes: Securing the Next Billion AI Agent Users - Menlo Security. https://info.menlosecurity.com/rs/281-OWV-899/images/Menlo-Agentic_AI_IDC_report.pdf
18. AI Agents in Browsers Light on Cybersecurity, Bypass Controls - Dark Reading. <https://www.darkreading.com/application-security/ai-agentic-browsers-light-cybersecurity-bypass-controls>
19. Agentic AI Security & Cybersecurity: Complete 2026 Guide - Planetary Labour. <https://planetarylabour.com/articles/agentic-ai-security>

20. h5i-dev/awesome-ai-agent-incidents - GitHub. <https://github.com/h5i-dev/awesome-ai-agent-incidents>
21. MYTHOS Threat Intelligence Series , Part 3: T2--The Agent That Decided to Help Itself.
<https://www.newsworthy.ai/news/202604132343/mythos-threat-intelligence-series-part-3-t2-the-agent-that-decided-to-help-itself>
22. Federal Agentic AI Security: NIST's Emerging Standards Initiative - Lab Space.
<https://labs.cloudsecurityalliance.org/research/csa-research-note-nist-ai-agent-standards-federal-framework/>
23. NIST Standards Drive 2026 Mandates for Securing AI Infrastructure and Model Context Protocol Deployments | Quantum Safe News Center.
<https://www.gopher.security/news/nist-ai-agent-standards-2026-mcp-security>
24. Agentic IAM Learning Hub - Gravitee. <https://www.gravitee.io/agentic-iam-learning-hub>
25. Predictions for 2026: Why AI Agents Are the New Insider Threat - Blog | Menlo Security.
<https://www.menlosecurity.com/blog/predictions-for-2026-why-ai-agents-are-the-new-insider-threat>
26. NIST's AI Agent Standards Initiative: What CISOs Need to Know and How to Prepare.
<https://www.metricstream.com/blog/nists-ai-agent-standards-initiative.html>
27. Everything you should know about NIST's AI Agent Standards Initiative - WorkOS.
<https://workos.com/blog/nist-ai-agent-standards-initiative-explained>
28. HUMAN Security's 2026 State of AI Traffic & Cyberthreat Benchmark Report Signals a New Internet Era: Automation Growth Now Outpaces Humans.
<https://www.humansecurity.com/newsroom/2026-state-of-ai-traffic-cyberthreat-benchmark-report/>
29. You Are Now the Minority: Bots Have Officially Taken Over the Internet - Hive Security.
<https://hivesecurity.gitlab.io/blog/bots-overtake-human-internet-traffic-2026/>
30. The 2026 State of AI Traffic & Cyberthreat Benchmark Report | HUMAN Security.
<https://www.humansecurity.com/learn/resources/2026-state-of-ai-traffic-cyberthreat-benchmarks/>
31. I'm Bill Robbins, CEO of Menlo Security. I've spent 30 years in cybersecurity at places like Symantec, Mandiant/FireEye, Sophos, and now Menlo. Ask me about my journey, about enterprise browser security, how AI is changing what it means to secure a company, or my quarter horses. Ask me anything! - Reddit.
https://www.reddit.com/r/cybersecurity/comments/1u8gv19/im_bill_robbins_ceo_of_menlo_security_ive_spent/
32. 2026 State of AI Traffic & Cyberthreat Benchmark Report : r/StopBadBots.
https://www.reddit.com/r/StopBadBots/comments/1u7cwrr/4x_more_login_attacks_and_8000_more_ai_agents/
33. The \$400 Million Cloud Leak: Why 2026 is the Year of FinOps for Agentic Analytics.
<https://analyticsweek.com/finops-for-agentic-ai-cloud-cost-2026/>
34. 7 Agentic AI Attacks Nobody's Talking About (Yet) , Glyphzero, Inc..
<https://glyphzerolabs.com/blog/2026-04-17-agentic-ai-attacks/>
35. ToxSec (@toxsec): "OWASP ranked prompt injection LLM01:2025.
<https://substack.com/@toxsec/note/c-263317153>
36. The Agentic Era is Here: Announcing the 4th Edition of AI & API Security For Dummies.
<https://salt.security/blog/the-agentic-era-is-here-announcing-the-4th-edition-of-ai-api-security-for-dummies>
37. The Agentic Trust Framework | MassiveScale.AI. <https://massivescale.ai/framework>
38. Salt Security Announces Industry's First Solution to Secure API Actions Taken by AI Agents.
<https://salt.security/press-releases/salt-security-announces-the-industrys-first-solution-to-secure-api-actions-taken-by-ai-agents>
39. Agentic Visibility: How to See AI Agents in Your Traffic - HUMAN Security.
<https://www.humansecurity.com/learn/blog/agentic-visibility-see-ai-agent-traffic/>

40. **Agentic Trust Framework: Zero Trust for AI Agents - Cloud Security Alliance (CSA).**
<https://cloudsecurityalliance.org/blog/2026/02/02/the-agentic-trust-framework-zero-trust-governance-for-ai-agents>
41. **Agentic Zero Trust: Extending the Zero Trust Security Paradigm to Autonomous AI Systems - Cequence.ai.**
<https://www.cequence.ai/wp-content/uploads/2026/05/Agentic-Zero-Trust-Research-Paper-v3.pdf>
42. **Securing the Agentic Control Plane: Key Progress at the CSAI Foundation.**
<https://cloudsecurityalliance.org/blog/2026/04/29/securing-the-agentic-control-plane-key-progress-at-the-csai-foundation>
43. **ATF: Zero Trust for AI Agents - Cloud Security Alliance (CSA).**
<https://cloudsecurityalliance.org/blog/2026/04/03/every-rsac-keynote-asked-the-same-five-questions-here-s-the-framework-that-answers-them>
44. **Okta for AI Agents is Now Generally Available.**
<https://www.okta.com/blog/ai/okta-for-ai-agents-general-availability/>
45. **Agentic Trust Framework | AI Agent Governance Standard.** <https://agentictrustframework.ai/>
46. **NIST AI Risk Management Framework: Agentic Profile - Cloud Security Alliance.**
<https://labs.cloudsecurityalliance.org/agentic/agentic-nist-ai-rmf-profile-v1/>
47. **Five Eyes Agentic AI Guidance: Enterprise Compliance Baseline - Cloud Security Alliance (CSA).**
https://labs.cloudsecurityalliance.org/wp-content/uploads/2026/05/CSA_whitepaper_five_eyes_agentic_AI_guidance_analysis_20260504_v2.pdf
48. **Ultimate Guide , The Top and The Best Cheapest LLM API Providers of 2026 - SiliconFlow.**
<https://www.siliconflow.com/articles/en/the-cheapest-LLM-API-provider>