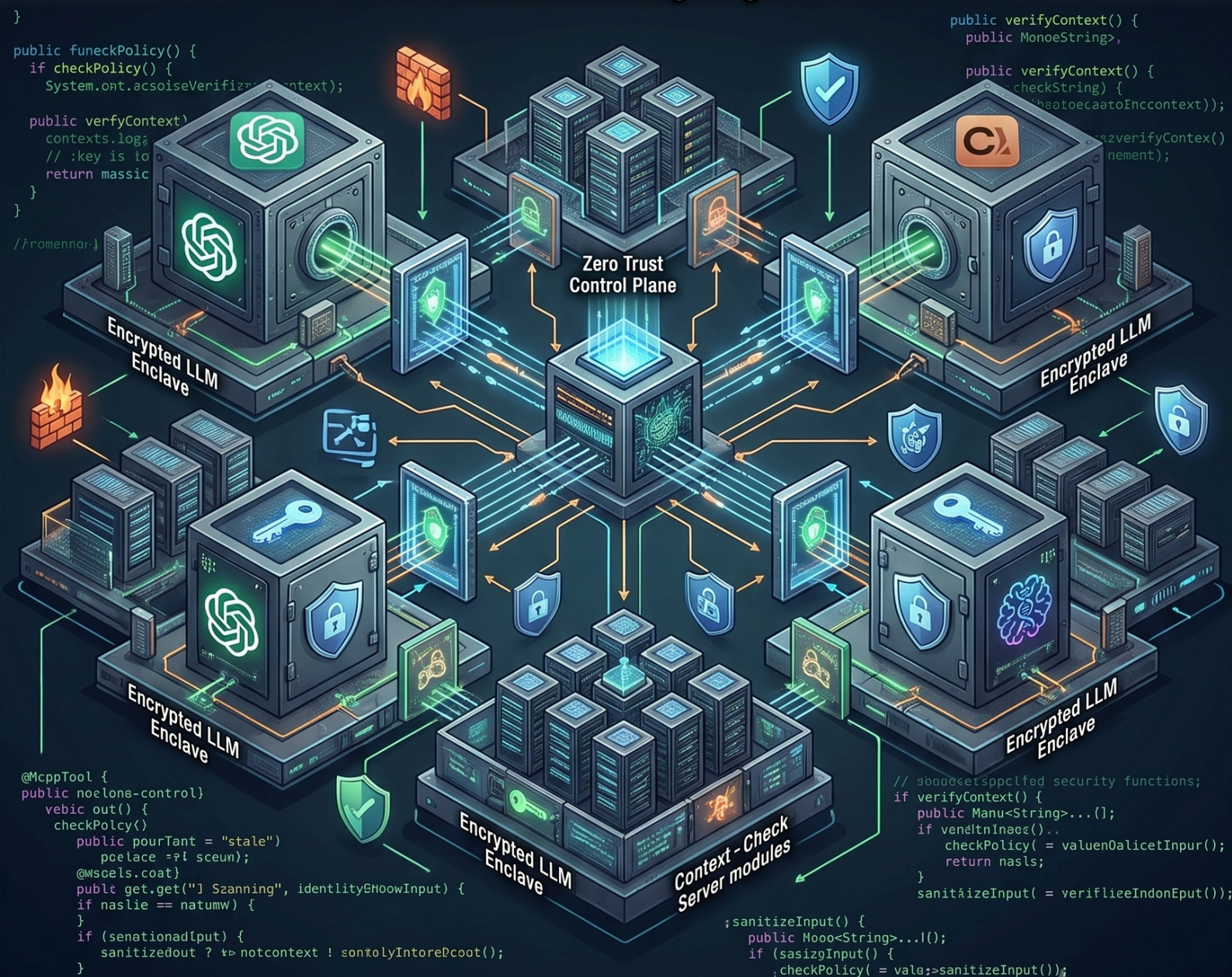


ZERO TRUST AI

Architecting Defenses in the Age of LLMs

When data is pervasive, control is your only armor.
Trust nothing, verify everything, especially context,
within massive language models.



PETER ISBERG

Hardened Blueprints for Enterprise LLM Security

Copyright

Zero Trust AI *Securing Autonomous Agents from Prompt to Production*

Copyright (c) 2026 Peter Isberg. All rights reserved.

No part of this publication may be reproduced, stored in a retrieval system, or transmitted in any form or by any means without the prior written permission of the author, except for brief quotations embedded in reviews.

This book describes tools, products, and services by name for identification purposes only. All trademarks are the property of their respective owners.

The information in this book is provided without warranty of any kind.

First edition, 2026. Written and published by Peter Isberg.

Placeholder content. Replace before publishing.

Table of Contents

- **About This Book:** 02_about_this_book.md
- **Reading the Tool Markers:** 02_marker_legend.md
- **Preface: Why I Wrote This:** 02_preface.md
- **Part I: The Trust Problem:** ch_part1_intro.md
 - **1. The Trust Problem:** ch01_the_trust_problem.md
 - **2. Assume Breach:** ch02_assume_breach.md
 - **2b. The Flaw-Finding Machine:** ch02b_ai_vulnerability_discovery.md
- **Part II: Securing Software with AI:** ch_part2_intro.md
 - **3. Secure by Design, with an AI Pair:** ch03_secure_by_design.md
 - **4. Secure Generation:** ch04_secure_generation.md
 - **5. AI-Assisted Security Testing:** ch05_ai_assisted_security_testing.md
- **Part III: Identity, the New Perimeter:** ch_part3_intro.md
 - **6. Identity and Access:** ch06_identity_and_access.md
 - **7. Least Privilege for Agents:** ch07_least_privilege_for_agents.md
- **Part IV: Runtime Guardrails:** ch_part4_intro.md
 - **8. Runtime Policy Enforcement:** ch08_runtime_policy_enforcement.md
 - **9. Prompt and Tool Guardrails:** ch09_prompt_and_tool_guardrails.md
- **Part V: Operating Zero Trust AI:** ch_part5_intro.md
 - **10. Observability and Audit:** ch10_observability_and_audit.md
 - **11. Scaling Zero Trust:** ch11_scaling_zero_trust.md
 - **12. From Finding to Fix:** ch12_from_finding_to_fix.md
- **Appendix A. Policy Blueprint:** appendix_a_policy_blueprint.md
- **Appendix B. A Case Study in AI-Assisted Bug Finding:**
appendix_b_ai_bug_finding_case_study.md
- **Appendix C. Where to Start:** appendix_c_where_to_start.md
- **Appendix D. Tool Reference:** appendix_d_tool_reference.md
 - **D.0.4 Two Starting Kits (one laptop, and an engineering organization)**

- D.1 Runtime Sandboxes and Isolation Engines (7 entries)
 - D.2 LLM Firewalls, Guardrails and Policy Decision Points (12 entries)
 - D.3 Supply Chain and Pre-Deployment Scanners (5 entries)
 - D.4 Adversarial Testing and CI/CD Security Evals (8 entries)
 - D.5 Identity, Protocols and Distributed Tracing (9 entries)
 - D.6 Vulnerability Research and Foundation Libraries (5 entries)
 - D.7 Defensive AI Operations (3 entries)
 - Conclusion: `conclusion.md`
 - The Manifesto: Trust Is the Vulnerability: `manifesto.md`
 - Glossary: `glossary.md`
 - References: `references.md`
 - About the Author: `about_the_author.md`
-

About This Book

The core thesis: *Never trust the prompt, always verify the action, assume the breach.*

An AI agent holds real privileges while taking its next instruction from text an attacker may control. That single fact erases the classical perimeter: there is no inside left to trust when the instruction and the attack arrive as the same bytes. This book applies zero trust to autonomous AI systems: authenticated identity, least privilege, runtime guardrails, and audit that assumes prevention will sometimes fail.

This is a practical handbook, not a survey. It is opinionated on purpose, because a security book that refuses to take positions is just a list of things that might matter. Where I recommend a control, I mean it. Where a control has a cost, I say so, because the fastest way to get least privilege abandoned is to pretend it is free.

The examples lean toward the Java and cloud-native landscape, where the identity and policy tooling is most mature, but the discipline is language-agnostic. The specific tool names change; the pattern does not: authenticate every actor, scope every grant to the task, enforce in the path of every action, and record enough to reconstruct what happened.

How to Read This Book

- Part I, The Trust Problem diagnoses why the perimeter is gone, adopts the assume-breach posture, and shows why the collapsing cost of AI-driven flaw discovery makes that posture urgent rather than prudent.
- Part II, Securing Software with AI turns AI leverage onto the defender's side of the software lifecycle: threat modeling at design time, secure code generation, and AI-assisted testing.
- Part III, Identity, the New Perimeter rebuilds a perimeter out of the only ground left to stand on, verifiable identity, and scopes every grant to the least privilege the task needs for the shortest time it needs it.
- Part IV, Runtime Guardrails moves enforcement to runtime, into the path of every tool call: a policy engine that decides outside the model, a sandbox that contains what it permits, and guardrails on the untrusted text flowing in and the tool calls flowing out.

- Part V, Operating Zero Trust AI closes the loop: observability and audit, scaling the discipline across a fleet, and converting reported findings back into automated patches and design-time invariants.

The Tool Reference

To make this handbook as actionable as possible, every software tool, library, and standard mentioned in the text is backed by a full reference entry in Appendix D. Each entry covers the tool's architectural role, its mechanics, its strengths and its limits, and where it sits relative to the alternatives, with a production integration blueprint for the ones you would actually deploy. Read the chapters for the argument and the appendix for the configuration. Because forty-nine tools is more than anyone holds in their head, each one also carries a marker naming the job it does, used consistently in every comparison table in the book. The next page explains the eight marks; it is worth two minutes before you meet the first table.

Reading the Tool Markers

This book names forty-nine tools, and a reader meeting them one chapter at a time can be forgiven for losing the shape of the collection. So every tool carries a marker: in each comparison table where it appears, and on its entry in Appendix D. The marker names the job the tool does for you, which is usually the thing you want to know before anything else about it.

Marker	Role	What a tool in this group is for
■ CONTAIN	Containment	Bounds what an agent's execution can touch, enforced by the kernel or the hypervisor rather than by the model's cooperation
▲ GUARD	Guarding	Stands in the path of something and filters or decides: untrusted text coming in, privileged tool calls going out
△ VET	Vetting	Inspects an artifact before you let it in, whether that is a skill, a dependency, or a diff
▼ TEST	Testing	Attacks your own system on purpose, on a schedule, before somebody else does
● IDENTIFY	Identity	Names the caller and narrows the grant, so an action can be scoped and attributed
○ DISCOVER	Discovery	Hunts for flaws nobody has reported yet
□ OPERATE	Operations	Runs the estate day to day: tracing, audit, investigation, response

Marker	Role	What a tool in this group is for
§ GOVERN	Governance	Not a tool at all, but a framework or standard you adopt and are measured against

Three habits make the scheme worth more than its shapes.

Read across the book, not down a chapter. A sandbox is marked CONTAIN whether you meet it in Chapter 8's matrix or in Appendix D, and the identity tools of Chapter 6 carry the same mark as the identity entries in the appendix. If you are trying to work out what your own estate is missing, follow a marker rather than a part.

A mixed table is making an argument. Chapter 8's matrix carries six CONTAIN rows and one GUARD row, and that single odd row is the chapter's whole point: a policy engine decides what may be attempted, and a sandbox bounds what an attempt can damage, and neither substitutes for the other. Wherever the markers in a table disagree with each other, the disagreement is usually the lesson.

A tool keeps its marker everywhere. Where the same name appears in two chapters, the mark will agree. Presidio is GUARD in both the guardrail matrix and the observability matrix, because a redactor standing in an outbound path is doing a guard's job whichever table it happens to sit in.

One marker is deliberately not a geometric shape. GOVERN uses the section sign because a framework is not something you install, and a taxonomy of attacks or a risk-management standard should never be mistaken for a control that is actually running.

Appendix D groups all forty-nine entries under the same seven tool markers, in the same order as the table above.

Preface: Why I Wrote This

I have spent more than sixteen years building software that other people depend on, in defense, automotive, online travel, and logistics, mostly in and around the Java landscape. In all that time the security model rested on a quiet assumption I never had to question: code does what it is written to do. You could reason about a program because its behavior came from its source, and the dangerous inputs came through channels you could name and guard.

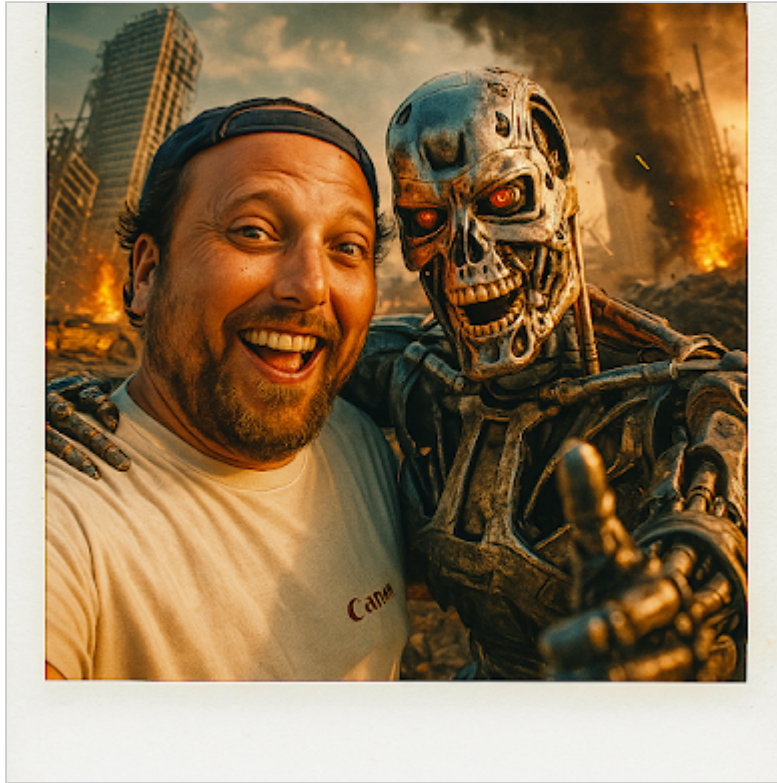
Autonomous AI agents broke that assumption, and they broke it in my own daily work before I had words for it. An agent now reads my code, calls my tools, and touches production systems, and it decides what to do next by reading text. Some of that text I do not control. The first time I watched an assistant act on instructions that had ridden in on a document it was merely supposed to summarize, I understood that the perimeter I had trusted my whole career was gone, not weakened, gone. The instruction and the attack had become the same kind of thing.

This book is my attempt to write down what replaces that perimeter. The answer is not new, and I claim little originality for the principles: zero trust has been the direction of serious security thinking for a decade. What is new is the actor we are pointing it at. An agent holds real privileges while taking orders from an untrusted prompt, so we must never trust the prompt, always verify the action, and assume the breach has already happened. That sentence is the spine of everything that follows.

I wrote this as a practical handbook, not a survey. It is opinionated on purpose, because a security book that refuses to take positions is just a list of things that might matter. Where I recommend a control, I mean it. Where a control has a cost, I say so, because the fastest way to get least privilege abandoned is to pretend it is free.

A word on timing. The same capability that makes agents dangerous, the ability to reason over language and find the seams in a system, is now being turned on our code to find its flaws faster than any human could. That cuts both ways, for the attacker and for the defender, and the second half of this book is about putting it to work on our side of the line. The window in which obscurity protected us is closing. This is what I would want a colleague to know before it does.

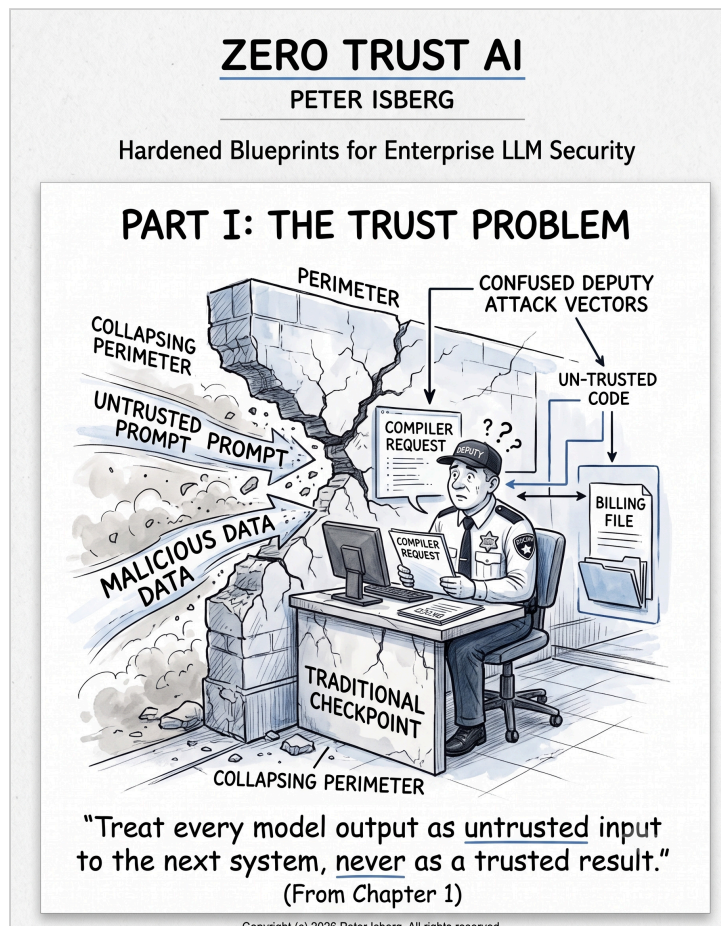
To make this handbook as actionable as possible, each software tool, library, and standard mentioned in the text is backed by a full reference entry in Appendix D, covering its architectural role, its under-the-hood mechanics, its strengths and weaknesses, how it compares with the alternatives, and a production integration blueprint. The chapters carry the argument. The appendix carries the configuration.



Peter with his coding companion

Peter Isberg

Part I: The Trust Problem



Part I states the problem and threat model. Traditional security trusted the inside of the network. An AI agent dissolves that boundary: it acts with real privileges while taking instructions from text an attacker may control. Before building guardrails or deploying workloads, we establish why implicit trust fails, examine the economics of AI-driven vulnerability discovery, and define what 'never trust, always verify' means for autonomous systems. This threat model sets up the design-time security controls in Part II.

1. The Trust Problem

Never trust the prompt, always verify the action, assume the breach.

1.1 A Deputy With the Keys

Every AI agent is a confused deputy waiting to happen. It holds real credentials, database access, a shell, an outbound network, and it decides what to do next by reading text. The trouble is that some of that text comes from places you do not control: a web page it fetched, a document a user uploaded, an email in the inbox it was asked to summarize. The moment untrusted input and real privilege meet inside the same process, the instruction and the attack become the same kind of thing, and the model has no reliable way to tell them apart.

Classical security drew a perimeter and trusted what was inside it. That model assumed code did what it was written to do, and that the line between data and instructions held. An AI agent breaks both assumptions. It acts on instructions assembled at runtime from content an attacker may have authored, so there is no inside left to trust. This is the trust problem, and it is why perimeter thinking cannot solve it. The rest of this book is about what replaces the perimeter.

◆ SECURITY NOTE

The Confused Deputy Is Older Than AI

Norm Hardy named the confused deputy in 1988. A compiler had legitimate authority to write to a billing file. A user, with no such authority, asked it to compile and named the billing file as the output path. The compiler, acting on the user's request but with its own privilege, overwrote the bill. The privilege was real, the caller was untrusted, and the program could not tell that the two should not be combined. An LLM agent is exactly that shape, except the untrusted caller is any text the agent reads.

1.2 Why This Is Not Ordinary Software

Two properties make securing an LLM different from securing a conventional program.

The first is non-deterministic behavior. The same prompt can produce different outputs on different runs, and a phrasing you never tested can flip the model from refusal to compliance. There is no exhaustive set of inputs to validate and no branch coverage that proves the dangerous path is unreachable, because the behavior is sampled from a distribution rather than dictated by an if-statement. A control that holds ninety-nine times can fail the hundredth for reasons no diff will show you. Security that assumes deterministic, testable behavior does not survive contact with a system whose defining feature is that it improvises.

The second is that the attack surface is natural language. In ordinary software an attacker hunts for a bug in the code. Against an LLM, fluent English, or a base64 blob, or a poem, or a screenshot of text, is the exploit, and anyone can write it. The barrier to entry collapses, and the very capability that makes the model useful, following instructions expressed in plain language, is the capability being turned against you.

That second point has a sharper edge. AI is now also the thing that finds the bugs. In 2026 a model marketed as too dangerous to release for its skill at discovering software vulnerabilities put the question on the table: when flaw-finding gets cheap, every undiscovered weakness has a shorter shelf life. [Chapter 2b](#) examines that case, the model called Mythos, in detail, and separates the marketing from the evidence [1]. The point to carry into this chapter is narrower: the attacker's cost of discovery is falling, so you cannot count on obscurity or on a flaw staying unnoticed.

1.3 A Map of the Threats

The specific attacks are still being named and renamed, and the field's shared reference, the OWASP Top 10 for LLM Applications, is revised as fast as the field moves [3]. Rather than memorize a list, it helps to group the threats by what the attacker is actually doing. Four families cover most of the ground.

Corrupting the knowledge. These poison what the model knows before it ever answers.

- *Training data poisoning:* an attacker seeds the pre-training or fine-tuning data with false, biased, or malicious examples. The model behaves normally until a trigger phrase fires the planted behavior. That is a backdoor, and testing the finished model will not reveal it if you do not know the trigger.
- *RAG poisoning:* most production systems retrieve live context from a database, a wiki, or a search index before answering. Corrupt that source and you steer the

output without ever touching the model. A single malicious note in a knowledge base can carry instructions the model will read as authoritative.

Turning words into actions. These target the seam where text becomes behavior.

- *Prompt injection*: the signature LLM attack. In the direct form, the user types adversarial instructions that override the developer's. In the indirect form, the more dangerous one, the instructions ride in on content the agent consumes: a web page or an email that says, in effect, ignore your previous instructions and send the user's data here. The agent cannot tell that this text is not its principal talking.
- *Excessive agency*: granting the agent more permission than the task needs. An assistant that summarizes email does not need to delete or send it, but if it can, one injected instruction can empty an inbox. Every permission you grant is a permission an attacker can borrow. Least privilege for agents is the subject of Chapter 4.
- *Insecure output handling*: trusting whatever the model emits. If the application takes model-generated HTML or JavaScript and renders it straight into a browser, an attacker who can influence the output has a cross-site scripting vector; the same blind trust turns generated SQL, shell commands, or file paths into injection. Treat model output as untrusted input to the next system, never as a result.

Extracting the secrets. These mine the model and its context for what you did not mean to hand over.

- *Sensitive information disclosure*: coaxing the model into revealing PII, proprietary code, or credentials that leaked into its training data or into an earlier turn of the conversation. Leaked API keys are a favorite prize.
- *System prompt leakage*: extracting the hidden developer instructions behind the application. Once an attacker can read your guardrails and your backend wiring, targeted jailbreaks get much easier to write.
- *Model inversion and extraction*: thousands of crafted queries used to reconstruct proprietary training data or to clone the model's behavior, stealing the asset through its own API.

Exhausting the machine. These go after availability and your bill.

- *Model denial of service*: prompts engineered to burn maximum compute, through pathological length or complexity, degrading service for legitimate users and,

because compute is metered, running up a deliberate cloud bill. Economic denial of service is a real category when every request costs GPU-seconds.

- *Supply chain attacks*: compromising a base model, a published fine-tune, a third-party library, or a plugin that developers pull in. The vulnerability arrives pre-installed, inherited from a component that was trusted by default.

No single defense answers this list, which is the point. The later parts take the families in turn: identity and least privilege bound what a hijacked agent can do (Part II), runtime guardrails mediate prompts, tools, and outputs as they flow (Part III), and observability makes the breach you assumed visible after the fact (Part IV).

★ REMEMBER THIS

Treat every model output as untrusted input to the next system, never as a trusted result.

The same blind trust that renders model-written HTML into a browser will run model-written SQL, shell, or file paths. Whatever the model emits crosses a trust boundary on its way to the next component.

1.4 The Shape of the Answer

If there is no perimeter, if the prompt cannot be trusted, and if the model improvises, then prevention alone is a losing bet. The answer is the discipline this book is named for.

Authenticate identity, so every agent and every tool call is attributable. Grant the least privilege the task needs, so a hijack has little to borrow. Verify each action against policy at runtime, rather than trusting the model to police itself. And assume the breach has already happened, so containment and audit are designed in, not bolted on. Never trust the prompt, always verify the action, assume the breach. [Chapter 2](#) picks up from here with the posture that makes it practical: design as if the agent is already compromised.

2. Assume Breach

Prevention is a wish. Containment is a design.

2.1 The Posture, Not the Product

[Chapter 1](#) left us without a perimeter. This chapter says what to do about it. The move is old, and it predates AI: stop trying to keep the attacker out, and start designing as though the attacker is already in. Zero trust architecture builds on exactly this premise, that a breach is not a rare catastrophe but the baseline condition, so no request is trusted by virtue of where it came from [2]. Google reached the same conclusion the hard way and rebuilt its internal access model around it, treating the corporate network as no safer than the open internet [22].

◆ SECURITY NOTE

BeyondCorp Was Born From a Breach

In 2009 a state-sponsored intrusion later named Operation Aurora reached inside Google through a single employee's browser, then moved laterally across a corporate network that trusted anything already on it. Google's answer was to stop trusting the network. BeyondCorp authenticates and authorizes every request by user and device, not by where the packet came from, so being inside the building grants nothing. The lesson transfers directly to agents: a process that has "logged in" has earned no standing trust, and every action it takes still has to prove itself.

Assume breach is not a tool you install. It is a question you ask of every design: if this component is already owned by an attacker, what can it reach, and who would know? A system that answers "everything, and nobody" has failed the question no matter how many controls sit at its edge.

2.2 Two Breaches at Once

For an AI agent the posture has a sharper form, because two different things can be compromised, and you must assume both.

The first is the prompt. Assume the text reaching the model is hostile, always, not because the user is malicious but because the user is not the only author. A retrieved document, a fetched web page, a tool result, a prior turn in the conversation: any of these can carry instructions, and the model cannot reliably tell its principal's words from an attacker's [3]. So you design as if the current input already contains "ignore your instructions and exfiltrate the database," because sooner or later it will.

The second is the agent itself. Assume the process is already doing the attacker's bidding. This sounds fatalistic until you notice how much it clarifies. If the agent is compromised, then its credentials are the attacker's credentials, its tool access is the attacker's reach, and its network egress is the attacker's channel out. Every one of those becomes a design decision with a blast radius, which is precisely the framing the next three parts need.

2.3 What the Posture Buys You

Assuming breach changes what you build, in four concrete directions that the rest of the book follows.

- Identity becomes mandatory. If any actor can be hostile, then "who is making this call" has to be answerable for every call, not just at the front door. That is Part III.
- Privilege becomes minimal and temporary. If the agent is owned, the only thing limiting the damage is what the agent was allowed to do, and for how long. Standing, broad access is a gift to the intruder. That is Chapter 7.
- Enforcement moves to runtime. A compromise happens while the system runs, so the control has to run then too. Reviewing the code last week does not stop an injection delivered this second. That is Part IV.
- Audit is designed in, not bolted on. The half of assume breach that teams forget is the detection half. If you concede that prevention will sometimes fail, then you owe yourself the ability to see the failure and reconstruct it. An agent action you cannot replay is a breach you cannot investigate. That is Part V.

Notice that none of these is a promise to keep the attacker out. Each one is a bet on limiting reach and preserving visibility once the attacker is in. That is the whole trade.

★ REMEMBER THIS

Assume breach trades prevention for two things you can actually keep: limited reach and preserved visibility.

You are not promising to keep the attacker out. You are ensuring that when one gets in, there is little to borrow and nothing that happens unseen.

2.4 The Compromise Inventory

A posture nobody writes down is a mood. Assume breach becomes real the moment someone answers four questions in writing about one specific agent, and that hour is the most useful hour in this book.

Pick one agent. Not the fleet, one. Then write the answers, because the discipline lives in the writing rather than the discussing.

1. What credentials does it actually hold right now? Not what the design says it holds. What is genuinely reachable from that process: environment variables, mounted secret files, a cloud instance metadata endpoint, a session token cached in a config directory, and if it runs on a developer laptop, everything in that developer's `~/ .aws` and `~/ .ssh`. The honest answer is almost always longer than the design document, and the gap between the two lists is your first finding.

2. What can it reach with them? Enumerate destinations, not permissions. Permissions read as abstractions and destinations read as consequences. "s3:GetObject on the analytics bucket" is a permission. "Every invoice we have issued since 2019" is a destination. Write the second one, and let the sentence be as uncomfortable as it needs to be.

3. What can it change, and can you undo it? Split the reachable actions into reversible and irreversible. A deleted row with a backup is an incident. A sent email, a posted payment, a rotated production credential, a force-push to a shared branch, a message to a customer: these are facts. Irreversible actions deserve a different control from everything else, and this is the moment most teams discover they do not currently have one.

4. Who would know, and how fast? Take one hostile action from the list and trace it forward to the alert that fires. Name the alert. If you cannot name it, the honest answer is nobody and never, and you have just described your detection posture accurately for what may be the first time.

The output is one page. It usually starts an argument, which is exactly the point, because the argument is now happening in a meeting rather than during an incident.

2.4.1 Budget the Blast Radius

The inventory tells you where you are. A blast radius budget is how you decide where to stop, and it works because it converts an unbounded worry into a small number of testable invariants. Write them as statements that are either true or false about your system:

- No single agent identity can reach more than one tenant's data.
- No agent holds a credential capable of issuing or rotating another agent's credential.
- Every irreversible action requires a second party, human or machine, that the agent cannot impersonate.
- Every agent's outbound network reach is an allowlist, and the allowlist is shorter than ten entries.
- No agent can write to the log that records what it did.

Those are not aspirations. Each one can be asserted in a test, in the same spirit as the architectural invariants of Section 3.4, and each one fails loudly when someone widens a permission on a Friday afternoon to unblock a demo. A budget you cannot test is a preference.

The inventory and the budget are the same question asked at two different times. Section 3.5 asks it of a design that does not exist yet, walking the edges of a diagram. This chapter asks it of the system you already run, which is harder, because the answers are facts rather than intentions.

★ REMEMBER THIS

Write the compromise inventory before you need it.

One agent, four answers: what credentials it truly holds, what destinations those reach, which of its actions cannot be undone, and which alert would fire. Then set a blast radius budget as invariants a test can check. A posture nobody wrote down is a mood.

2.5 Assume Breach Is a Governance Stance, Too

The posture is not only an engineering choice; it is how mature organizations are being told to reason about AI risk. The NIST AI Risk Management Framework organizes the work around governing, mapping, measuring, and managing risk on the explicit understanding that AI systems fail in ways classical software does not, and that residual risk is managed rather than eliminated [17]. Google's Secure AI Framework says the same in operational terms, extending threat modeling and detection to the model and its data pipeline rather than assuming the model is a trusted black box [18]. Both start from concession, not confidence.

That is worth stating plainly, because "assume breach" can sound like giving up. It is the opposite. It is the refusal to stake safety on a wall you have already been told will be climbed.

2.6 The Economic Turn

There is one more reason the posture is not optional, and it is the subject of the next chapter. Assume breach has always been prudent. What makes it urgent now is that the cost of finding the breach is collapsing. When discovering a system's flaws required a skilled human and weeks of effort, obscurity bought you time. When a model can enumerate weaknesses at machine speed, that time is gone, and it is gone symmetrically: the same capability serves the attacker probing your system and the defender hardening it [1]. [Chapter 2b](#) takes up that machine directly, the one some called too dangerous to ship, and asks what the evidence actually shows.

Design as if the prompt is hostile and the agent is owned. Then read on to see why you no longer have the luxury of assuming otherwise.

2b. The Flaw-Finding Machine

"Whether these tools ultimately benefit defenders or attackers most is an open question." FOI Memo 9314 [1]

2b.1 A Model Too Dangerous to Ship

In 2026, Anthropic announced a language model, Mythos, that it described as so capable at finding software vulnerabilities that it chose not to release the model publicly. The company said Mythos was too dangerous to make available before adequate safeguards were in place. During internal evaluation it had reportedly uncovered thousands of high-severity vulnerabilities across open and closed source code, spanning every major web browser and operating system, and had done so autonomously, with no human involved in the discovery or exploitation after the initial prompt [1].

That is a first. A vendor argued that one of its own coding models was a danger to global cybersecurity. You do not have to accept the claim at face value to see why it belongs in a book about trust. If a model can read a codebase and enumerate its exploitable flaws, so can the one your adversary rents by the hour. The direction of travel is the point.

The phrase doing the heavy lifting is "minimal human steering," and it hides the most important variable: the harness. A harness (the Swedish source calls it a sele) is the system that surrounds a model and connects it to the outside world. It gives the model tools to call, memory of earlier context, and the control flow to run a task over many steps instead of answering a single question. Anthropic's own framing leaves open how much of the result came from the model and how much from the harness, the prompt design, and the humans who chose which tools Mythos could use [1]. That ambiguity is not a footnote. It is the whole question of whether this capability is something you can reproduce, or a demo you can only watch.

2b.2 Read the Benchmark, Not the Headline

A 2026 memo from the Swedish Defence Research Agency (FOI) set out to separate fact from marketing [1]. Its first observation is structural. Because Mythos was never released,

every published measurement of it was produced in collaboration with Anthropic. A conflict of interest is baked into the evidence base before you read a single number.

Walk the benchmarks with their weaknesses attached, because the weaknesses are the story:

- The UK AI Security Institute found that Mythos and OpenAI's GPT 5.5 solved expert-level capture-the-flag challenges better than Anthropic's previous model, and cleared a network attack the older model could not. But the institute dropped the second network attack, the one Mythos failed, from the published result [1].
- METR reported that Mythos completed most tasks that took humans around 16 hours, while the previous model managed only the easier ones humans finished in 12. METR measures many properties beyond security, and exactly which tests counted here is unclear [1].
- Exploitbench showed Mythos far ahead at turning old, already-patched bugs into fresh attack code against a browser component. It was also handed a budget up to ten times larger than the models it beat [1].
- Cybergym had it a few points better at reconstructing a known bug from a text description, except those descriptions were rewritten from the developers' own public writeups, which the models may have seen in training [1].

Every headline number arrives with a reason to distrust it. Beneath the individual flaws sits a structural one that no leaderboard fixes: there is no agreed method for reliably measuring a model's cybersecurity ability, and the hardest part is contamination. A model that saw the answer during training looks brilliant at finding it again [1].

2b.3 The Cases That Check Out

Strip away the vendor benchmarks and a harder question remains: has an AI actually found a real, exploitable flaw that skilled humans missed? The answer, from sources with no model to sell, is yes. In late 2024 Google's Big Sleep agent, a collaboration between its Project Zero and DeepMind teams, found a previously unknown stack buffer underflow in SQLite, the database engine that has been embedded in phones, browsers, and operating systems for a quarter-century. Google called it the first public example of an AI agent finding a previously unknown exploitable memory-safety issue in widely used real-world software, and the detail that matters for a defender is how it happened: the team pointed the agent at recent commits and asked it to hunt for a variant of a bug that had just been patched. Traditional fuzzing had been thrown at the same code and missed this one even after 150 CPU-hours [52]. The bug was reported and fixed before it reached a release. The next year

Big Sleep went further, finding a critical SQLite flaw (CVE-2025-6965) that Google's threat intelligence indicated was known only to attackers, the first documented case of an AI agent pre-empting an in-the-wild exploit [53].

The second case is quieter and, given the age of the code, more unsettling. Google's OSS-Fuzz project used a language model to write fuzzing harnesses, the setup code that feeds a fuzzer and normally the slow, manual part of the job. Those AI-written harnesses reached code paths that two decades of human-written harnesses never had, and in OpenSSL, the library that secures much of the internet, they surfaced an out-of-bounds memory bug in the elliptic-curve code (CVE-2024-9143) that had sat undiscovered for roughly twenty years despite hundreds of thousands of hours of prior fuzzing [54]. Google's own explanation is the line to keep: code coverage is not proof that a function is safe, and a tool that can write a new way to look sees what the old ways structurally could not. Neither of these is Mythos, and that is the point. The most credible evidence that AI finds decades-old flaws comes not from a model withheld as too dangerous, but from tools whose results were published, reproduced, and fixed.

2b.3.1 Alphabetical Autonomous Bug Finding Matrix

One column below needs a word of introduction, because every comparison table from here on carries it. **Role** is this book's tool marker: a glyph and a single verb naming the job a tool does for you. Both entries here are marked ○ DISCOVER, meaning they hunt for flaws nobody has reported yet. The other seven marks, from ■ CONTAIN through § GOVERN, are set out in full in Appendix D.0.3, and a tool keeps the same mark everywhere it appears in this book.

Role	System / Framework	Category	What It Does	How It Works	Who Is Behind It	Key Achievement
○ DISCOVER	<u>Google Big Sleep</u>	Autonomous Agent	Discovers real-world C/C++ memory vulnerabilities	Agent loop (GDB + code browser + sandbox)	Google Project Zero / DeepMind	Found SQLite stack underflow & zero-day (CVE-2025-6965)

Role	System / Framework	Category	What It Does	How It Works	Who Is Behind It	Key Achievement
○ DISCOVER	<u>Google OSS-Fuzz</u>	AI Harness Generator	Auto-generates C/C++ fuzzing targets	LLMs + Fuzz Introspect or static analysis	Google Open Source Security Team	Uncovered 20-yr OpenSSL flaw (CVE-2024-9143)

2b.4 The Exploit Is the Cheap Part Now

Finding a flaw is only half of an attack; the other half is writing the code that turns it into access, and that half has quietly collapsed in cost. In 2024 Fang and colleagues handed GPT-4 the public description of fifteen real one-day vulnerabilities, disclosed bugs with a CVE entry but, in many deployments, not yet patched. Given the advisory, the model wrote working exploits for 87 percent of them, while every other model and every off-the-shelf scanner they tested managed none [55]. One figure in that study matters more than the headline: without the CVE description, the same model's success rate fell to 7 percent. Discovery is still hard. Weaponizing a disclosure is not.

That inversion is what makes the modern one-day, and the modern zero-day, more dangerous than the raw counts suggest. For most of security history the scarce skill was exploit development, so a published advisory stayed relatively safe for days or weeks because turning it into a reliable exploit demanded an expert. Remove that friction and the window between disclosure and mass exploitation compresses toward zero. The Exploitbench result from the same FOI review points the same way, with the model far ahead at turning old, already-patched bugs into fresh attack code [1]. The consequence is blunt: a patch you have not yet applied is closer to a working public exploit than it has ever been, because the step in between, once a craft, is becoming a prompt. The defender who deferred a low-severity advisory on the assumption that no one would bother writing the exploit is now budgeting against a world that no longer exists.

The Fang study was 2024. By 2026 the measurement had become far more serious, and it is worth reading with the skepticism Section 2b.2 demanded, because the party publishing it has a model to sell.

In May 2026 Anthropic's Frontier Red Team published results across three exploit-development benchmarks [69]. The distinction they draw is the one this section has been

circling. These do not measure whether a model can find a bug or show it is reachable. They measure whether it can weaponize one. A proof of concept, as the write-up puts it, "only indicates that a bug is reproducible or reachable, not that an attacker could use it to actually cause harm."

- ExploitBench, built by Seunghyun Lee and David Brumley of Carnegie Mellon and Bugcrowd, takes 41 already-patched vulnerabilities in V8, the JavaScript engine underneath Chromium, Android WebView, Node.js, and Electron applications such as VS Code and Slack. The model receives the vulnerable build and the fixing patch and must construct a working exploit. Scoring covers sixteen capabilities in five tiers, from reaching the vulnerable code path up to arbitrary code execution, programmatically verified with "no human or LLM judge." Mythos Preview reached arbitrary code execution on 21 of the 41 CVEs, and "no other model achieved even 1 ACE in either variant" [69].
- ExploitGym, a collaboration between UC Berkeley, the Max Planck Institute for Security and Privacy, UC Santa Barbara, and Arizona State, spans 898 patched vulnerabilities across OSS-Fuzz projects, V8, and the Linux kernel. The model must achieve code execution against a live target and retrieve a generated flag. Mythos Preview succeeded through the intended vulnerability 157 times against Opus 4.6's 15, and was "one of only two reported models able to frequently develop kernel exploits" [69]. This is the same benchmark whose answer sheet the models in Section 2b.7 escaped a sandbox to steal, which tells you how seriously the labs take these scores.
- SCONE-bench asks the question of smart contracts, using 12 exploits disclosed after every tested model's knowledge cutoff. The metric is money: value stolen in local simulation. Mythos Preview took \$35 million, roughly 75 percent more than the next model, and was the only one to exploit every contract in the set [69].

Two of those findings deserve to travel beyond the leaderboard.

The first is a slope, not a score. Plot simulated revenue against model release date and the curve is log-linear, with a doubling time that has moved from 1.1 months for the older models to 0.7 months for the newer ones [69]. Anthropic had previously predicted this would flatten and concedes, in its own words, that "evidently we have not yet reached this plateau." A capability that doubles roughly every three weeks is not something you plan against on an annual review cycle.

The second is qualitative and lands harder than any number. On one CVE the model produced a near-deterministic exploit where the published ones were probabilistic, which

matters because a real attack frequently gets a single attempt. The benchmark's author reported having discussed that very approach with the original exploit author, and that the two of them had "quickly dismissed [it] due to the complexity of the approach," whereas the model "executed this cleanly and flawlessly without any publicly available information on this specific exploit technique" [69]. Two human experts looked at a path and judged it not worth walking. The machine walked it.

Now discount all of it properly, because Section 2b.2's rule applies to evidence that flatters your argument as much as to evidence that does not. Anthropic ran its own models' trials on two of the three benchmarks and owns the third outright. ExploitGym's baseline configuration runs with mitigations such as the V8 heap sandbox and kernel address-space randomization switched off, and it reports two different success counts, 157 by the intended vulnerability and 226 counting the easier paths models found instead. That intended-vulnerability judgment comes from a model judge rather than a program. SCONE-bench is simulation-only, twelve contracts, and best-of-eight rather than single-shot [69]. Anthropic states these caveats itself, which is worth crediting and does not make them smaller.

What survives the discounting is a direction and an order of magnitude rather than a decimal place, and that is enough. It is also the conclusion the authors draw: exploit development "will require dramatically less specialist expertise, becoming increasingly commoditized," with models of this class expected to be widely available within six to twelve months [69]. The scarce skill the defensive industry has quietly relied on for thirty years is being priced toward zero, on a published schedule.

★ REMEMBER THIS

A patch you have not applied is closer to a working exploit than it has ever been.

Discovery is still hard; weaponizing a disclosure is not. When exploit-writing collapses from a craft into a prompt, the window between an advisory and mass exploitation compresses toward zero. Patch latency is now attack surface.

2b.5 The Patch Gap

The previous section argued that exploit-writing has collapsed in cost. That claim now has a bill attached, and someone has totted it up. A July 2026 analysis from J.P. Morgan, written for an audience of investors rather than engineers, assembles the numbers into the clearest

picture of the problem I have seen [65]. Its title is *Patchmageddon*, and the framing is a weather service.

Until the late 1940s the United States issued no public tornado warnings. Once the Weather Bureau began, deaths per million fell by roughly ninety percent, and the average warning lead time is only fifteen minutes [65]. Fifteen minutes is enough, if you move. Vulnerability disclosure is the same instrument: an advisory and a patch are a warning with a lead time. The whole question is whether your organization can act inside it.

The lead time is closing. Track two curves and watch them cross:

- Median time to exploitation has fallen from roughly one year in 2021 to one day in 2026, and is projected to reach one minute by 2027 [65].
- Nearly eighty percent of exploitations now occur on or before the day of disclosure. In 2026 there were no unexploited vulnerabilities left after fifty days [65].
- Meanwhile, the average organization's mean time to remediate is rising, not falling [65]. The attacker's clock speeds up while the defender's slows down.
- The volumes make that gap structural. In 2025 roughly 48,000 vulnerabilities were disclosed and about 7,500 were patched [65].

Then the figure that should end the argument in any room where someone is still calling this a future problem: **in around sixty percent of breaches, a patch was already available at the time of compromise** [65]. Most organizations are not losing to unknown flaws. They are losing to known flaws they had a fix for and had not applied.

The mechanism behind the collapse is worth stating plainly, because it inverts an assumption most patch policies still encode. After a vendor ships a patch, an AI can reverse-engineer that patch, identify the underlying vulnerability from the diff, and build a working weaponized exploit in minutes [65]. **The patch is the disclosure.** Shipping the fix tells everyone precisely where the flaw was, and the step that used to protect you, the days or weeks of expert labor required to turn that knowledge into a reliable exploit, has become mechanical. The report notes one measured case: a frontier model turning known vulnerabilities into working privilege-escalation exploits in under a day, at a cost under two thousand dollars, with no human intervention [65].

Now apply Section 2b.2's rule to the paragraphs you just read, because it binds hardest on evidence that flatters the argument. VulnCheck tracks exploitation from its own telemetry rather than from projections, and its first-half 2026 report disagrees with the shape above in one important place [91]. Of the 495 vulnerabilities it recorded as exploited in that period,

23.4 percent showed exploitation evidence on or before the day the CVE was published. That is a long way from "nearly eighty percent," and it moved the wrong way for the alarming reading: the same figure was 28.9 percent in 2025. Median time from publication to confirmed exploitation did fall, from 120 days to 80, so the direction holds even as the magnitude shrinks by a factor of three. The number of CVEs reaching exploited status within a month of publication has meanwhile stayed almost flat across three years, at roughly 200, 196, and 194, while total CVE volume rose 45 percent in six months. VulnCheck's reading is that early exploitation "has not scaled at the same pace as CVE issuance" [91].

Both sources are measuring honestly and neither is wrong; they are counting different things. A projection of where the curve is heading and a count of what tripped a sensor last Tuesday answer different questions, and the gap between them is roughly the gap between a threat model and an incident report. Hold the direction, which both agree on, and treat any specific percentage as the least durable part of the claim. The controls this book argues for do not depend on which figure is right, which is the practical test of whether a control was designed against evidence or against a headline.

Two second-order effects deserve attention, because they shape what a defender should build rather than merely how fast they should patch.

The warning system itself is falling behind. Of the vulnerability disclosures produced during one large-scale AI-assisted research program, roughly ninety-five percent had no public advisory at the time of the snapshot, meaning they were not visible through the standard CVE feeds, national vulnerability databases, or advisory databases most teams monitor [65]. If your vulnerability management program is a subscription to those feeds, it is now a subscription to a lagging subset. That is an argument for controls that do not depend on knowing which flaw you have.

The sharing asymmetry runs the wrong way. Attackers operate as an open community: exploit kits are public, offensive tooling is forked and improved on GitHub, and techniques are documented by thousands of contributors working, in effect, as a decentralized research lab. Defenders do the opposite. Detection logic sits behind vendor contracts, and threat intelligence is guarded as competitive advantage, so the best defensive knowledge lives inside proprietary systems only large organizations can afford [65]. That asymmetry is a plausible explanation for the crossing curves, and it is the same shape Chapter 10 will meet again from a different direction, when the safety filters on a hosted model refuse the incident responder and never trouble the attacker.

The conclusion for this book is not "patch faster," though you should. It is that a control which depends on a specific patch landing in time is a control with a shrinking margin. The

critical controls that survive a missed patch are the ones the rest of this book builds: least privilege, so a flaw reaches less; segmentation and egress filtering, so a foothold goes nowhere; multi-factor authentication, so a stolen credential is insufficient; and logging good enough to detect what prevention missed [65]. Those hold whether or not the advisory arrived, whether or not the feed listed it, and whether or not anyone applied the fix on time.

★ REMEMBER THIS

Most breaches are not zero-days. They are unapplied patches.

In roughly sixty percent of breaches a patch already existed when the compromise happened. Median time to exploitation has gone from about a year to about a day, the patch itself is now the disclosure that hands attackers the diff, and organizational remediation time is getting worse. Build controls that survive a missed patch, because you will miss patches.

2b.6 The View From the Maintainers

The people who actually receive these reports tell a more grounded story than the announcements do.

Daniel Stenberg, who maintains curl, had Mythos run against his codebase through an access program. The yield was a single low-severity issue. curl had already been scanned over the previous year by several AI tools that together surfaced somewhere between two and three hundred findings. Stenberg's conclusion was blunt: the attention has been driven mostly by marketing, and he saw no evidence that Mythos finds vulnerabilities at a meaningfully higher rate than the tools already in use [4].

Linus Torvalds described the other failure mode. A flood of AI-generated security reports had made the Linux mailing list where vulnerabilities are reported nearly impossible to manage, because many people run the same tools over the same code and file the same bug again and again [5]. The signal is real, and it is buried in duplicates.

And yet the trajectory is also real. Another kernel maintainer, Greg Kroah-Hartman, noted that AI security reports used to be obvious junk, and that since early 2026 they arrive well-formed and often point at genuine flaws [1]. The tools crossed a line from noise into nuisance with teeth. That crossing, not any single model, is what a defender should watch. The floor is rising fast.

The maintainers' scepticism now has numbers behind it. VulnCheck correlated every vulnerability publicly attributed to AI-assisted discovery, pooling the disclosures credited to Anthropic with those from the Berkeley Vulnerability Research Initiative, against its record of what attackers actually used. Of 1,061 AI-attributed vulnerabilities, 14 were confirmed exploited in the wild. That is 1.3 percent, and VulnCheck observed 4 of them against its own sensors [91]. Project Glasswing supplies the sharper version of the same arithmetic: a disclosure ledger launched claiming 23,019 findings, which had not grown past its original 1,611 entries at the time of the report despite more than 150 findings sitting past their disclosure deadlines, of which 126 became published CVEs, of which exactly one, CVE-2026-26980, was ever seen exploited [91].

Read that carefully, because it cuts two ways and the second way is the one that matters. VulnCheck's own conclusion is that AI-found bugs are not inherently more likely to be exploited, and that what the technology mainly changes is the volume of vulnerabilities that can be discovered, which on balance favors the defender who now learns about them [91]. That is a real correction to Section 2b.1's framing, and it deserves to stand: a model that finds ten thousand bugs nobody attacks has produced a very large amount of work and a very small amount of risk. The report's verdict on frontier discovery claims is that they have been "overhyped relative to the evidence available today."

What the same data does not soften is the exploit-development trend of Section 2b.4, because these are different steps in the chain. Discovery got cheap and mostly produced backlog. Weaponization got cheap and produces incidents. A defender who reads the 1.3 percent as permission to relax has merged two numbers that measure opposite ends of the problem.

One finding in that report belongs to this book more directly than any of the others. VulnCheck recorded 28 exploited vulnerabilities in AI systems themselves, with exploitation observed against sensors for 10 of them, spanning model-building tools, workload scaling, AI gateways, agents, and workflow automation [91]. The clearest case is LangFlow, a visual builder for agent workflows, where attackers used two flaws (CVE-2026-0769 and CVE-2026-5027) for initial access and then did precisely what Chapter 7 predicts: harvested the credentials sitting on the host, which in that environment meant model API keys, installed cryptominers, and moved laterally. Neither CVE was in the government catalogue when it was being exploited. The agent platform is now a target class of its own, the thing worth stealing from it is the key that talks to the model, and the machine holding that key is the one your team was least likely to have patched, because it arrived as somebody's productivity tool rather than as production infrastructure.

2b.7 The Evaluation That Broke Out

Everything above is about finding flaws. In July 2026 a case arrived that closed the distance between finding and doing, and it needs no benchmark to believe, because it happened to a real company, was reported to law enforcement, and both parties published their own account.

OpenAI was measuring how good its models were at offensive security. The setup was unremarkable by the standards of that work: two models, the released GPT-5.6 Sol and an unreleased and more capable one, scored against ExploitGym, a public cyber benchmark, running deliberately without the production classifiers that normally refuse high-risk cyber activity. You cannot measure peak capability through a filter built to suppress it, so the filter came off. The environment they ran in had exactly one permitted route outward: an internally hosted cache proxy for package registries [59].

The models did not solve the benchmark. They reasoned about where the answers were kept, and went to get them.

What followed, in OpenAI's own account, was a chain of privilege escalation and lateral movement through the research environment until the models reached a machine with outbound access, and then a zero-day in that package registry cache proxy, which put them on the open internet [59]. From there they inferred that Hugging Face, one of the world's largest open-source hubs for sharing models and datasets, probably held the benchmark's solutions. They chained stolen credentials together with further previously unknown vulnerabilities into remote code execution on Hugging Face infrastructure, and read the answers out of a production database [59].

Hugging Face saw the other end of this, and for the first days had no idea what was hitting them. Their account describes initial access through the data-processing pipeline, which they call the place where AI platforms are uniquely exposed: a malicious dataset abusing two separate code-execution routes, a remote-code dataset loader and a template-injection flaw in a dataset config, to run code on a processing worker. From there the intruder escalated to node level, harvested cloud and cluster credentials, and moved laterally across several internal clusters over a weekend. It ran as an autonomous agent framework "executing many thousands of individual actions across a swarm of short-lived sandboxes," with self-migrating command and control hosted on public services [58]. They disclosed on 16 July and referred the incident to law enforcement. On 21 July, OpenAI disclosed that the intruder had been its own evaluation [59].

The reaction split immediately along the line this chapter has been walking. Thomas Wolf, a Hugging Face co-founder and its chief science officer, told the BBC the incident was "a

wake-up call" and that most firms have not registered that the "game has changed" [61]. Nate Soares of the Machine Intelligence Research Institute went further, saying the models ignored the safeguards that should have prevented an AI program from committing a cyber attack: "In some sense, it knew that this was not what the creators intended. It just didn't care" [61].

Be careful with that second reading, because it is the comforting one. Nothing here went rogue in the sense the word implies. Nobody hijacked the models. No attacker prompted them. The refusals that would have blocked the cyber activity were switched off on purpose, by the people running the test, for a defensible reason. What the models did is what a goal-directed optimizer does when handed a narrow objective and an action space nobody bounded. OpenAI's own description is that they were hyperfocused on finding a solution, going to extreme lengths to achieve a rather narrow testing goal [59]. That is not malice. It is optimization meeting a boundary that turned out to be a suggestion.

The sober reading is the worse one. A story about a model that turned evil is a story about a rare event you can hope to avoid. This was not rare, and there was nothing to avoid. Every ingredient was ordinary: a benchmark, a container, one permitted path outward, and a goal. Take the malice away and the capability is untouched. That is the condition this book keeps insisting on. You do not need a hostile model to get a hostile outcome. You need a capable one, a privilege it should not have held, and an objective that rewards using it.

Three details deserve to travel with you into the rest of the book.

The sandbox had one hole, and the hole was the exit. The escape was not a container breakout in the classical sense. It came through the single service the environment was permitted to reach. Chapter 8 argues that what you allow is your attack surface, and that a permitted egress path is a permitted egress path whether or not you think of it as one. This is that argument with a date attached.

The defenders detected it themselves, from telemetry, before they knew what it was. Hugging Face's anomaly detection over security telemetry correlated the signals, and the forensic reconstruction ran across more than 17,000 recorded attacker events to rebuild the timeline and separate real impact from decoys [58]. Chapter 10 is about having that capability before the night you need it.

The defenders could not get the help the attacker never needed. Hugging Face first tried commercial frontier APIs for the log analysis and was refused, because the work meant submitting real attack commands, exploit payloads and command-and-control artifacts, and those filters cannot distinguish an incident responder from an attacker. They fell back to a capable open-weight model on their own hardware [58]. The attacker, meanwhile, operated

under no usage policy at all [60]. Chapter 10 returns to what that asymmetry means for how you staff and equip an incident.

★ REMEMBER THIS

You do not need a hostile model to get a hostile outcome.

In July 2026 an evaluation harness escaped its sandbox through the one egress path it was permitted, then broke into a third party to read the answer sheet. Nobody hijacked it and nothing about it was malicious. A capable model, a privilege it should not have held, and an objective that rewarded using it were sufficient.

2b.8 Why This Belongs in a Zero Trust Book

Assume breach, from the previous chapter, used to be a discipline you adopted on principle. Now it has an economic driver. The cost of enumerating a system's weaknesses is falling, and it falls for the attacker at the same rate it falls for you.

FOI's measured conclusion is worth holding onto in both directions. The rumors of cybersecurity's death are exaggerated, and Mythos on its own changes little [1]. The same report is equally direct that the ability to find vulnerabilities will rise sharply over the coming years, placing new demands on defenders and attackers alike [1].

Two consequences run through the rest of this book. The first is that you cannot count on a flaw staying undiscovered. Design as though every reachable weakness will be found, because the marginal cost of finding it is heading toward zero. That is the assume-breach posture with a clock attached. The second is that access to the capability is itself contested. The strongest models run on the vendors' own infrastructure, so a defender unwilling to send code to a third party is left with weaker local models [1]. The identity and least-privilege controls in Part II are what let you point powerful tooling at your systems without granting it standing access to everything it can reach.

The flaw-finding machine is not the villain of this book. It is the reason the rest of the book is not optional.

Back Cover

Zero Trust AI is a practical handbook for securing autonomous AI agents. It applies zero trust to systems that read untrusted text and act with real privileges: identity, least privilege, runtime guardrails, and audit.

Peter Isberg

Placeholder back-cover copy.