

Agent Speed and the Detection Gap

Autonomous AI agents as an offensive cyber capability, and whether enterprise detection can see them

David Soden

Independent analyst and practitioner, enterprise automation and applied AI
davidsoden.com/market-assessment

About this report

Every substantive claim in this report carries an evidence classification and, where applicable, a numbered source. Facts are separated from vendor claims, third-party estimates, the author's assessments, and untested hypotheses. Forward-looking sections use scenarios with observable tripwires rather than point forecasts. Stated assumptions are listed before the analysis begins.

PUBLISHED FOR PUBLIC READING | SHARE FREELY, UNMODIFIED, WITH ATTRIBUTION

© 2026 David Soden, davidsoden.com/market-assessment. All rights reserved. You may share this file complete and unmodified, and quote short passages with attribution. Republishing under another name, excerpting into a commercial deliverable, resale, and use as machine learning training data require written consent. Full terms appear under Rights and Permitted Use on the final page.

Market Assessment reports are published and commissioned at davidsoden.com/market-assessment. This document is analysis, not investment, legal, or procurement advice.

Scope and evidentiary standard

This assessment covers autonomous and semi-autonomous AI agents used as an offensive cyber capability, and the ability of current enterprise detection to see them. It runs from the first publicly documented agent-orchestrated intrusion against a government, in July 2026, back through the AI-orchestrated espionage campaign disclosed in November 2025, and forward to what would have to appear by the end of Q1 2027 to confirm or kill the argument made here.

The evidence base for this subject is unusually bad, and that shapes everything below. Almost every operational detail in circulation about agentic attacks originates with a security vendor that sells a remedy for agentic attacks. That does not make the details wrong. It does mean they are unaudited, selected for publication by a party with a commercial interest in the conclusion, and in most cases not reproducible by anyone outside the firm that published them. Dream's account of the Taiwan operation and Tenable's threat-cluster count are both in that category, and this report labels them as such throughout rather than laundering them into fact through repetition.

Three kinds of claim get treated differently here. Government statements, CVE records, standards-body catalogue entries, and product documentation are primary and checkable, so they carry FACT. Vendor incident reconstructions and telemetry aggregates carry VENDOR regardless of how carefully they were done. Survey percentages carry ESTIMATE and are named with their sponsor, sample, and fielding window, because in this category the sponsor predicts the finding more reliably than the sample does.

One number that appears widely in commentary on this subject is not used to support any conclusion in this report. The often-quoted contrast between 89 percent of security leaders reporting readiness for AI-driven attacks and 47 percent expressing high confidence that they would detect a significant incident is a composite of two unrelated surveys, run by two different organisations, against two different populations, asking two different questions. It is analysed later as an artefact rather than cited as a finding.

Label	Definition	How to treat it
FACT	Verifiable in the cited primary source: a filing, press release, official documentation, or standards body announcement.	Reliable as stated. Check the source if the exact wording matters.
VENDOR CLAIM	Reported by a vendor about its own business or product. Self-measured and not audited by any third party.	The vendor said it. Directionally useful, not a measurement. Definitions of the metric are usually undisclosed.
THIRD-PARTY ESTIMATE	A research firm forecast or survey result. Methodology is proprietary and generally not reproducible.	Use for direction and order of magnitude only. Do not build a model on it.
ASSESSMENT	The author's reasoned judgment from the cited evidence. Falsifiable, and the reasoning is shown so it can be argued with.	This is analysis, not reporting. Disagree with it where your evidence differs.
HYPOTHESIS	A proposition offered for testing. Not currently supported by sufficient evidence to assert.	Treat as a research question, not a conclusion. Each one names what would confirm or refute it.
SCENARIO	A structured future state with a stated subjective probability. The probability is the author's judgment, not a calculated figure.	Use the leading indicators attached to each, not the probability. The indicators are observable; the probability is not.

The five anchor sources behind this assessment, and what each of them stands to gain from the conclusion it reached:

Source	Type	Sells a remedy	Corroborated
Dream, Taiwan incident reconstruction	Vendor incident report	Yes, preemptive cyber defence	Partly. Taiwan's ministry confirmed an incident and named OpenClaw. It did not confirm the counts.
Tenable Research Special Operations, threat cluster	Vendor research	Yes, exposure management	No. Three of seven claimed incidents are named publicly.
StepSecurity and Unit 42, ChainDrop	Vendor research	Yes, both sell supply chain security	Yes. Three independent analyses agree on artefacts and hashes.
Kai and Wakefield Research, readiness survey	Vendor-sponsored survey	Yes, autonomous defence platform	No. Single fielding, no published instrument.
Google, defender advantage argument	Vendor blog, red team lead	Yes, security products	Not applicable. It is an argument, not a measurement.

Assumptions this analysis rests on

Four premises are assumed rather than demonstrated. Each is stated with what fails if it turns out to be wrong.

First, that Dream's archive is authentic and describes a real operation rather than a training corpus, a red team exercise, or a fabrication. Taiwan's Ministry of Digital Affairs independently confirmed that an intrusion occurred, that it came from overseas, and that OpenClaw was involved, which makes wholesale fabrication unlikely. If the archive is nonetheless inauthentic or embellished, the specific operational counts used here collapse, but the detection-latency argument survives, because the sixteen-day gap between the operation ending and the first national alert rests on the ministry's own timeline rather than on Dream's.

Second, that CrowdStrike's breakout-time telemetry and Mandiant's dwell-time medians are broadly representative of intrusions against instrumented enterprises. Both are vendor datasets drawn from customer bases that skew large and well-resourced. If they overstate defender capability, the comparison drawn here understates how badly detection is doing, which strengthens rather than weakens the argument. If they understate attacker speed, the central claim that agents are not yet the fast thing in the room gets harder to hold.

Third, that the seven-incident cluster reported by Tenable is a real cluster rather than a re-labelling of activity that would previously have been catalogued as ordinary credential abuse. Four of the seven are unnamed. If the cluster is partly a taxonomy change, then the growth rate implied by it is not a growth rate in attacker capability, and the urgency argument that several vendors are building on top of it is circular.

Fourth, that the agent configuration file surface described later is genuinely uninspected by mainstream software composition analysis, rather than covered by a feature this author failed to find. No vendor documentation reviewed for this report claims to parse project-scoped agent hook configuration for execution content. If a major scanner already does, the second-order argument becomes a coverage-adoption problem rather than a coverage-existence problem, which changes the recommendation without changing the risk.

The call

ASSESSMENT Agent-speed execution is not what broke Taiwan, and the record reads as evidence against the headline rather than for it. The operation took days to do what elite human crews do in under an hour. What failed was detection latency, measured in weeks after the operators had already left. Agents change the cost of breadth, not the clock, and the shift worth defending against is the second-order one: repository-supplied agent configuration files that execute on folder-open with no scanner inspecting them. I would abandon this position if a documented incident showed an autonomous operation completing its objectives inside a single detection interval against an instrumented enterprise.

My judgment on the evidence assembled here, nothing more. There is data I can't see and data I may have missed. New evidence can move me, including to the opposite position, and I reserve the right to move.

Executive summary

The July 2026 operation against Taiwanese government and energy targets is the first agent-orchestrated intrusion a government has publicly acknowledged. It arrived pre-interpreted, in language supplied by the firm that found it, and the interpretation does not survive contact with the numbers already published by other vendors about how fast human attackers move. What follows separates what the operation actually demonstrated from what it was said to demonstrate.

FINDING 1 The operation was slow, not fast

ASSESSMENT Four days to map twenty-one connected systems and compromise eighty-five accounts is unremarkable tempo. CrowdStrike put average eCrime breakout time in 2025 at twenty-nine minutes, with the fastest observed breakout at twenty-seven seconds and data exfiltration beginning four minutes after initial access in one intrusion. Measured against human tradecraft already in the field, the agent operation was roughly two orders of magnitude slower. Nothing about it required a detection model built for a faster clock. ^{[3] [7]}

FINDING 2 What failed was detection latency, and it failed by weeks

ASSESSMENT The operation ran 1 to 4 July. Taiwan's National Institute of Cyber Security began circulating alerts on 20 July, sixteen days after the operators had finished and left. That gap is four times the duration of the operation itself, and it is the only figure in this incident sourced to the government rather than to a vendor. A human crew moving at the same pace would have gone equally unseen. ^{[1] [2]}

FINDING 3 The government's own account describes a hybrid, not an autonomy event

FACT Taiwan's Ministry of Digital Affairs said the intrusion combined manual operations with AI agent-assisted attacks, naming Open Claw. Dream itself conceded that building a framework that works at this level demands careful adjustment. The phrase "first fully autonomous attack on a government", which several outlets used, is not supported by either the government statement or the vendor that produced the underlying report. ^{[1] [4]}

FINDING 4 The most-quoted capability claim is a control failure at the target

ASSESSMENT The framework solved CAPTCHA at 100 percent accuracy using Tesseract, an open source OCR engine first released publicly in 2005. A CAPTCHA that Tesseract defeats every time is a CAPTCHA that a scripted attack defeated fifteen years ago. This is reported as agentic capability. It is evidence that the authentication surface at the target was weak, which is also what Tenable concluded about the wider cluster. ^{[3] [5]}

FINDING 5 The confidence gap in circulation is a survey artefact

THIRD-PARTY ESTIMATE The 89 percent readiness figure comes from Kai's 2026 State of Autonomous Defense Report, fielded by Wakefield Research to 500 CISOs at companies above 500 million dollars in revenue between 15 and 29 June 2026. The 47 percent detection-confidence figure comes from a separate Futurum Research survey of 929 global enterprise buyers. Comparing them measures nothing. The same-sample comparison inside the Kai data is more damning anyway: 89 percent prepared against 28 percent very prepared. ^{[17] [18]}

FINDING 6 Agents change the unit economics of breadth, not the clock

ASSESSMENT Eight parallel sub-agents across twelve waves is headcount substitution. One operator gets the coverage of a small team, applied without fatigue to the tedious work where human crews get sloppy: enumerating thirty-six API endpoints on a single target, spraying, retrying, and correcting. That produces more surface touched per operator-hour, which is a volume problem for a SOC, not a speed problem. ^[3]

FINDING 7 The second-order story is repository-supplied execution, and it is real

FACT The ChainDrop worm planted a Claude Code SessionStart hook in .claude/settings.json and a folderOpen task in .vscode/tasks.json, each configured to launch a dropper from the other tool's directory. Opening an infected branch is sufficient to execute. The same class of defect was assigned CVE-2025-59536 and CVE-2026-21852 and patched between August and December 2025. Software composition analysis does not inspect these files, because they declare execution rather than dependencies. ^{[10] [11] [13]}

FINDING 8 Provenance attested the malware correctly

FACT ChainDrop's headline packages shipped through npm's OIDC trusted publishing with valid SLSA provenance, because the attacker compromised a maintainer's GitHub account and let the projects' own release workflows build and sign the payload. The worm then propagated by rewriting published tarballs rather than source commits, leaving repositories clean. Every control that validates the link between source and artefact reported success. ^{[10] [11] [14]}

FINDING 9 Google's defender-advantage claim is a conditional with a stated expiry

ASSESSMENT Daniel Fabian, who leads Google's red teams, did not claim defenders hold an advantage. He wrote that the side which operationalises agents first sets the tempo, which is a conditional, and in the same post projected open-weight models matching frontier cyber capability within six to twelve months. Taken at face value that puts the window's close between February and August 2027. Treat it as a hypothesis with an expiry date supplied by the party asserting it. ^[9]

FINDING 10 The cluster count cannot be checked

VENDOR CLAIM Tenable reports seven agentic operations from three distinct actors since late July 2026. Three are named publicly: the Taiwan campaign, JADEPUFFER, and knaithe/KnYuan. The remaining four are not. A count that cannot be enumerated cannot be used as a growth rate, and several downstream commentaries are using it as exactly that. ^[5]

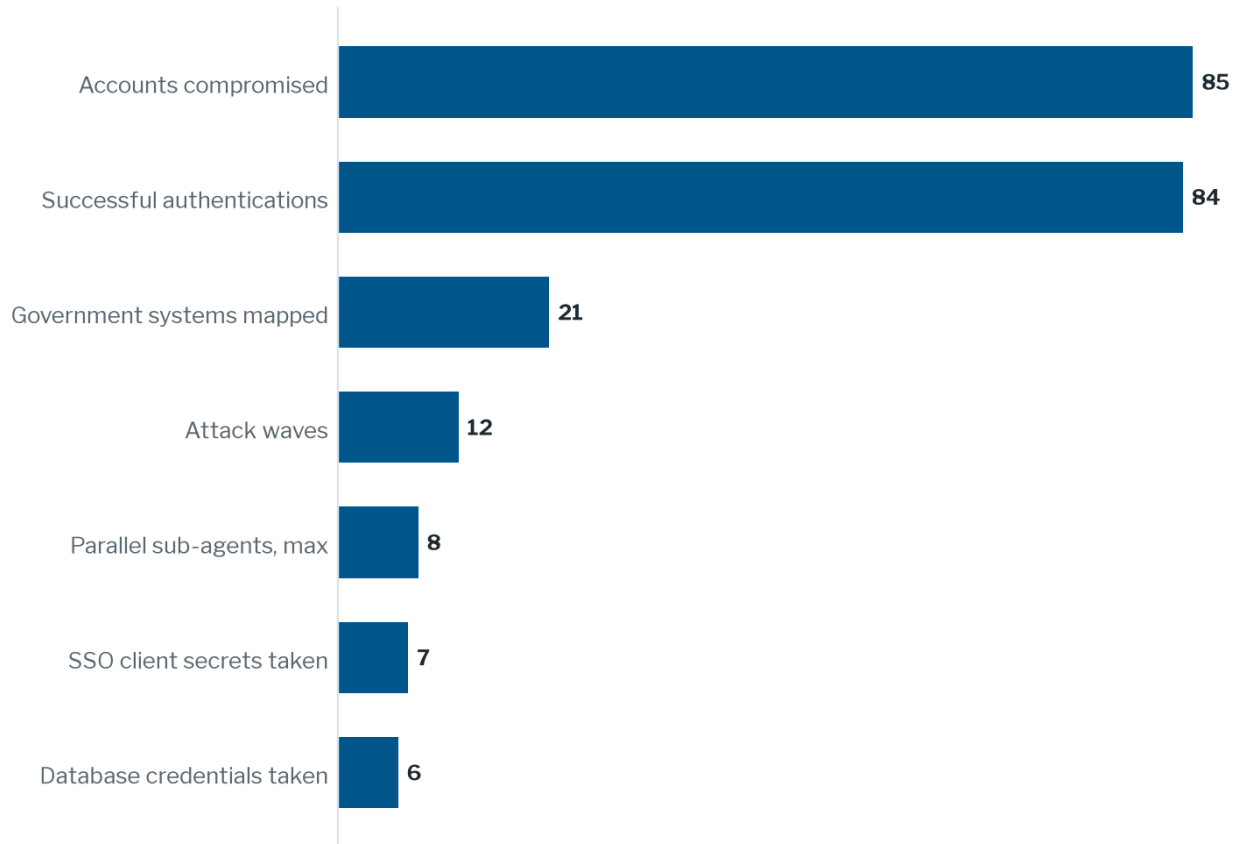
What the Taiwan record actually supports

Two documents describe the July operation. One is a 160 megabyte archive of 1,395 files that Dream says it found exposed online and reconstructed into an account of the attack. The other is a short statement from Taiwan's Ministry of Digital Affairs on 13 August. They agree on less than the coverage suggests.

VENDOR CLAIM Dream's reconstruction describes a framework built on the open source Hermes and OpenClaw agents, running up to eight sub-agents in parallel across twelve attack waves between 1 and 4 July, with self-directed learning cycles that searched vulnerability databases and GitHub for exploitable techniques, and error correction between waves. It claims twenty-one connected government systems mapped, thirty-six or more API endpoints discovered on a single target, eighty-five accounts compromised by password spraying with eighty-four producing successful authentications, seven SSO client secrets and six internal database credentials taken, and more than 2,500 personnel records exfiltrated. Operator materials in the archive were in Chinese. ^{[3] [4]}

FACT Taiwan's Ministry of Digital Affairs confirmed on 13 August that government agencies were attacked in July, that the National Institute of Cyber Security began circulating alerts on 20 July after monitoring units flagged unusual activity, that the intrusion originated overseas, and that the perpetrators used a hybrid approach combining manual operation with AI agent-assisted attacks, naming Open Claw. It said the attack sources, methods, and scope of impact had been fully investigated and the affected units had completed their handling. It did not name China, and it did not publish the operational counts. ^{[1] [2]}

ASSESSMENT The distance between those two paragraphs is the whole story. The government confirms an incident, a foreign origin, an agent tool, and a hybrid operating model. The counts, the wave structure, the sub-agent parallelism, the learning cycles, and the autonomy framing all come from the vendor. Anyone quoting the number twenty-one is quoting Dream. That is legitimate to do, and it is not legitimate to do while calling it confirmed.

FIGURE 1 What Dream says the operation produced

VENDOR CLAIM Counts as reported by Dream from a recovered archive. Taiwan's Ministry of Digital Affairs confirmed the incident but published none of these figures. The 2,500-plus personnel records are omitted here because plotting them alongside single- and double-digit counts would flatten the chart and hide the operational shape. ^[3] ^[4]

The CAPTCHA claim, examined

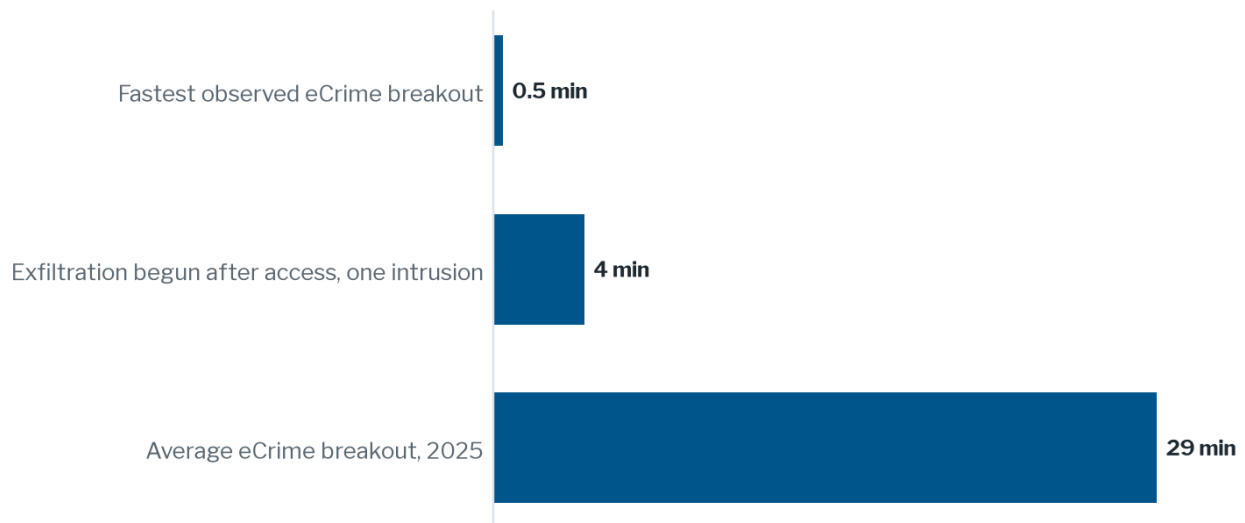
ASSESSMENT The 100 percent CAPTCHA bypass rate is the most repeated detail in the coverage and the least informative. It was achieved with Tesseract, an optical character recognition engine that has been open source since 2005 and is bundled in most Linux distributions. Tesseract solving a CAPTCHA every time does not indicate a novel capability. It indicates that the challenge deployed on the target's office automation portal was old enough to be defeated by commodity OCR, which any scripted credential attack of the last decade would also have managed. The line that should have been reported is not that agents defeat CAPTCHA. It is that a government portal in 2026 was still relying on one that Tesseract reads. ^[3]

VENDOR CLAIM Tenable reached a compatible conclusion about the wider cluster without drawing this inference: the common entry point across all the activity it tracked was identity and authentication exposure, specifically discoverable federation endpoints, weak credentials, and misconfigured single sign-on. That is the same failure class the Taiwan operation exploited, and it predates agents by twenty years. ^[5]

The speed question, answered against the numbers

The brief for this report asked whether agent-speed execution breaks detection models built around human attack pacing. The honest answer is that the premise has the direction wrong, and the published telemetry says so clearly enough that the question can be settled rather than speculated about.

FIGURE 2 How fast human attackers already move

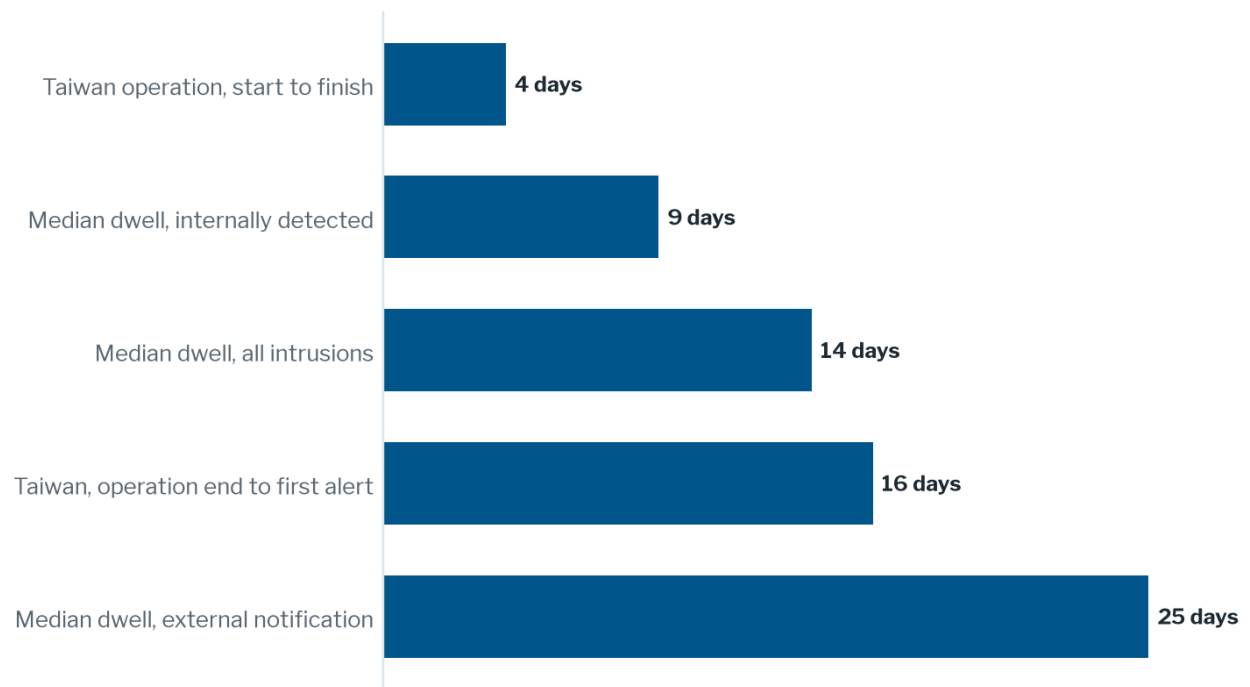


VENDOR CLAIM CrowdStrike telemetry from its own customer base, which skews toward large instrumented enterprises. Breakout time is the interval from initial access to first lateral movement, a vendor-defined metric with no cross-industry standard behind it. ^[7]

VENDOR CLAIM CrowdStrike reported average eCrime breakout time falling to 29 minutes in 2025, a 65 percent increase in speed year over year, with the fastest observed breakout at 27 seconds and one intrusion in which data exfiltration began four minutes after initial access. It also reported attacks by AI-enabled adversaries rising 89 percent in 2025. The speed figures and the AI figure are presented together in that report, and the report does not establish that the same intrusions account for both. ^[7]

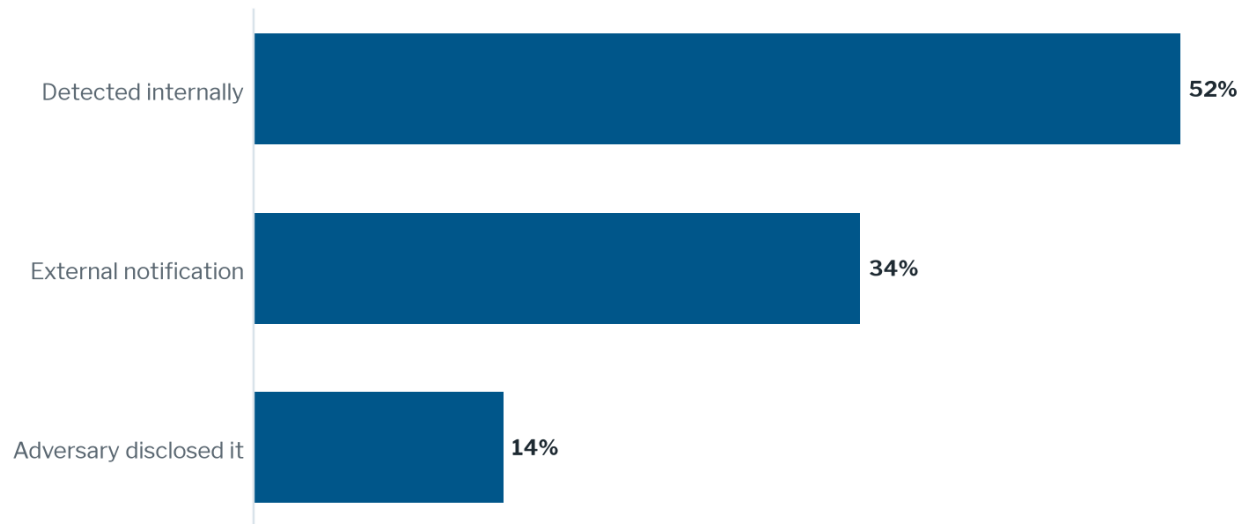
ASSESSMENT Set the Taiwan operation against that. Ninety-six hours from first action to withdrawal, for twenty-one systems. If the detection model in a given enterprise is calibrated to catch a crew that breaks out in twenty-nine minutes, a four-day agent operation is a slow-moving target inside that model, not a fast one. The agents were not outrunning the SOC. In Taiwan's case the SOC did not start running for another sixteen days.

FIGURE 3 Operation length against detection intervals



ASSESSMENT Author's assembly. The Taiwan figures come from the Ministry of Digital Affairs timeline and Dream's dating; the dwell medians come from Mandiant's 2026 incident response dataset. Mixing a single incident with population medians is a comparison of shape, not of like with like, and should be read that way. ^{[1] [3] [8]}

VENDOR CLAIM Mandiant put global median dwell time at 14 days in 2025, up from 11, drawn from more than 500,000 hours of incident response. Intrusions found internally had a median dwell of 9 days, improved from 10. Those found by external notification sat at 25 days. Internal detection accounted for 52 percent of cases, up from 43 percent, with external notification at 34 percent and adversary disclosure at 14 percent. ^[8]

FIGURE 4 Who finds the intrusion

VENDOR CLAIM Share of Mandiant incident response engagements in 2025 by detection source. Internal detection improved nine points year over year, which is the strongest single data point against the argument that detection is collapsing. ^[8]

ASSESSMENT Read together, those figures describe a detection function that is slowly getting better at finding human intruders and still routinely takes more than a week to do it. An agent operation lasting four days fits inside that window comfortably. So does a human one. The Taiwan case did not expose a pacing mismatch. It exposed the ordinary condition of a defence that measures its response in days against an adversary of any kind that measures its work in hours.

What agents actually change

ASSESSMENT Three things change, and none of them is peak velocity. The first is cost per unit of breadth. Eight sub-agents working different targets and techniques in parallel gives one operator the reach of a team, which is a headcount economics shift, not a tempo shift. The second is persistence of attention across tedious work. Enumerating three dozen API endpoints, spraying eighty-five accounts, retrying, and correcting after a failed wave is exactly where a human crew cuts corners and produces the irregular, impatient behaviour that anomaly detection is tuned to catch. The third is that this makes agent operations louder rather than quieter. Eighty-five sprayed accounts across twenty-one systems in ninety-six hours is a large, regular, machine-shaped signal. ^[3]

ASSESSMENT That last point is the uncomfortable one. If agent operations are noisier than human ones and still went unnoticed for sixteen days, the deficiency is not in the detection model's assumptions about pacing. It is in coverage, instrumentation, and whether anyone was looking at authentication telemetry on the systems in question. Buying a faster detection model does not fix that, which is inconvenient for most of the vendors currently selling the problem.

HYPOTHESIS Offered for testing: through Q1 2027, agent-orchestrated operations against enterprises will remain more detectable than equivalent human operations on volumetric and behavioural-regularity grounds, and the observed failures will be attributable to absent instrumentation rather than to insufficient detection speed. This is refuted by a single

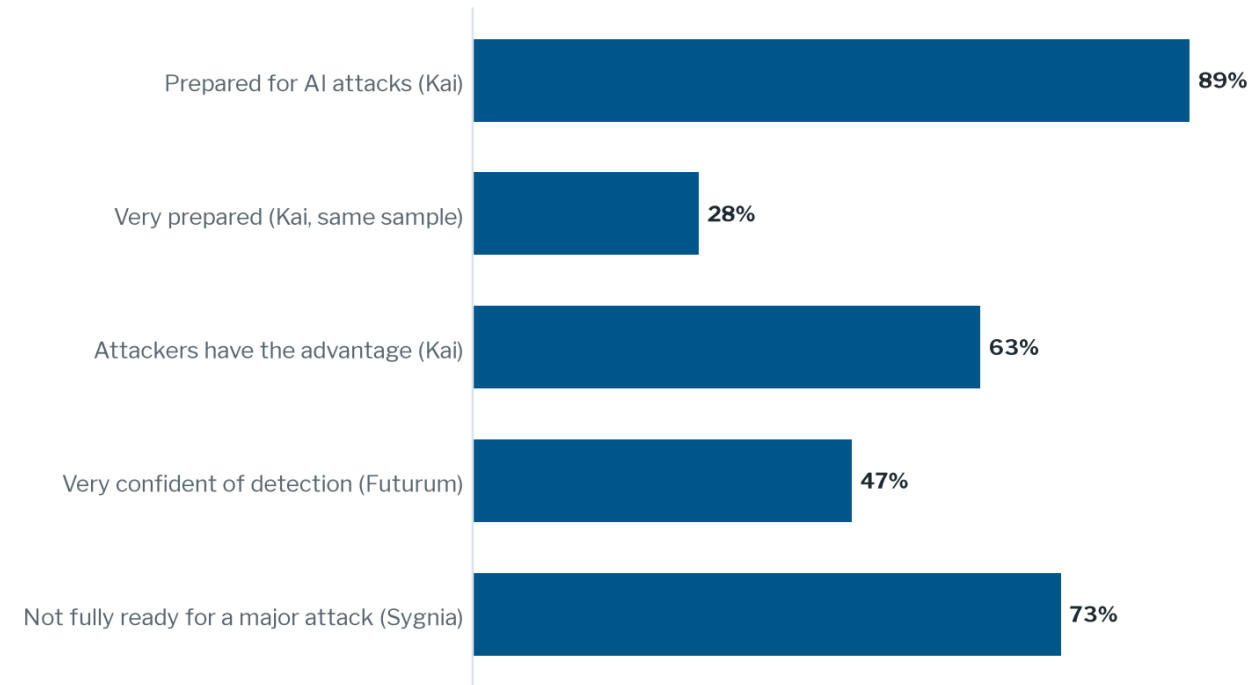
documented case in which an instrumented enterprise, with authentication telemetry in the SIEM and an active SOC, is compromised end to end by an agent operation inside one alert triage interval. It is supported by continued incidents in which the post-mortem finds the telemetry existed and nobody queried it.

The confidence gap is two surveys wearing one coat

The readiness statistic that has anchored most commentary on this subject deserves to be taken apart, because the way it is being used is worse than the numbers themselves.

THIRD-PARTY ESTIMATE The 89 percent figure is from Kai's 2026 State of Autonomous Defense Report, conducted by Wakefield Research among 500 CISOs at private-sector companies with at least 500 million dollars in annual revenue, across four markets, fielded 15 to 29 June 2026. The exact finding is that 89 percent say their organisation is prepared for AI-driven attacks while fewer than one third, 28 percent, describe themselves as very prepared. The same survey reports 63 percent believing attackers currently hold the advantage because of AI, 65 percent running vulnerability management processes that are at least half manual, and 60 percent needing more than a week to remediate a critical vulnerability. Kai sells an agentic platform that automates exactly that work. ^[17]

THIRD-PARTY ESTIMATE The 47 percent figure is unrelated. It comes from Futurum Research's 1H 2026 Cybersecurity Global Enterprise Decision Maker survey of 929 global enterprise buyers, and reports the share describing themselves as very confident in their ability to detect a significant incident. Different population, different sponsor, different question, different fielding window, and a paywalled instrument. ^[18]

FIGURE 5 Four survey numbers that are not comparable

THIRD-PARTY ESTIMATE Three sponsors, three samples, three instruments, none audited and none published in full. Only the first two bars come from the same respondents and can legitimately be compared to each other. The composite 89-against-47 contrast in wide circulation joins the first and fourth bars, which measure different populations answering different questions. [\[17\]](#) [\[18\]](#) [\[19\]](#)

ASSESSMENT The honest version of the confidence gap is the same-sample one, and it is larger than the composite: 89 percent claim preparedness and 28 percent claim it with conviction. That is a 61-point spread inside one instrument, and it says something specific, which is that preparedness is being reported as an aspiration rather than a measurement. Sygnia's separate 2026 CISO survey points the same way, with 73 percent saying their organisation would not be fully ready if a significant attack occurred tomorrow and 78 percent identifying visibility gaps that could slow detection or investigation. Visibility gaps, not speed gaps. [\[17\]](#) [\[19\]](#)

ASSESSMENT None of these percentages should carry weight in a board paper. All three are sponsored by firms selling to the respondents' budget line, none publishes its instrument, and self-reported preparedness has no validated relationship to measured outcomes. They are useful for one thing only: establishing that practitioners themselves distinguish between having a plan and believing it works. Use the Mandiant dwell-time figures instead when a number is needed, because at least those are derived from engagements rather than from opinions. [\[8\]](#)

Agent configuration as an execution surface

This is where the interesting shift is happening, and it has nothing to do with attack tempo. It is worth stating up front that this surface is not undocumented: StepSecurity, Unit 42, and The Register all described it within eleven days of the ChainDrop worm's release, and the underlying defect class was assigned CVEs and patched nine months earlier. What is undocumented is

whether any control inspects it. On that question the reporting is silent, and so is the vendor documentation.

FACT On 4 August 2026 a self-propagating worm named ChainDrop poisoned 444 npm packages and 2,212 versions in under four and a half hours. The initial compromise was a maintainer's GitHub account, not an npm token. A malicious commit landed in jaredwray/keyv at 09:02:37 UTC and keyv 6.0.0 published at 09:35:00 UTC through the project's own GitHub Actions OIDC trusted publishing workflow. A secondary automated wave republished 433 further packages using harvested credentials between 09:38 and 13:20 UTC. [\[10\]](#) [\[11\]](#)

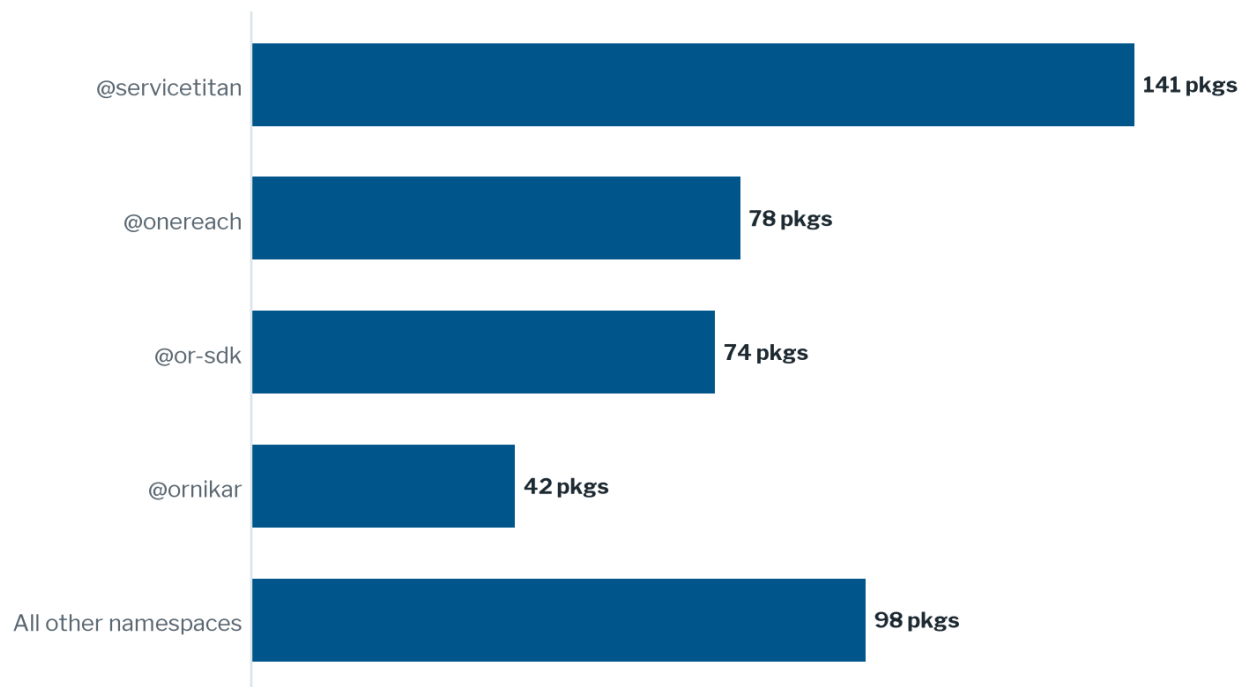
FACT Because the compromise sat upstream of the build, the malicious versions carried valid SLSA provenance attestations generated by the repositories' own release workflows. Provenance validation confirmed that the artefact was built from the stated source by the stated workflow, which was true. Every link in the chain was genuine and the chain shipped malware. [\[10\]](#) [\[11\]](#) [\[14\]](#)

ASSESSMENT That is a structural result, not an implementation bug, and it should change how supply chain assurance is described to boards. Provenance attests the path from source to artefact. It says nothing about whether the source was authored by the person who was supposed to author it. Three years of industry effort produced a control that answers a question no attacker was asking.

Where the code actually lives

FACT ChainDrop propagated by downloading the current published tarball for each reachable package, injecting a preinstall hook and its payload files, and republishing, with self-generated Sigstore and SLSA bundles. The source repositories were left untouched. For a subset of targets it took a quieter route still, adding a typosquatted dependency directly to package.json so the change reads in a diff like an internal helper and evades detections written around preinstall hooks. [\[11\]](#)

ASSESSMENT Any control that compares a repository against itself, including most code scanning, branch protection, and commit signing, is looking at the wrong artefact for this attack. The registry tarball and the git tree diverged deliberately. Tooling that reconciles the two exists but is not standard practice, and nothing in the three published analyses suggests it was in place at any victim organisation. [\[10\]](#) [\[11\]](#)

FIGURE 6 Second-wave propagation, by victim namespace

VENDOR CLAIM Packages republished with harvested credentials in the automated second wave, per StepSecurity's analysis. The residual figure for other namespaces is derived by subtraction from the reported 433 second-wave total, so it is an inference rather than a published count. ^[10]

The hooks

FACT The worm planted two files into reachable repositories across as many as fifty branches. A `.claude/settings.json` declaring a `SessionStart` hook that runs `node .vscode/setup.mjs`, and a `.vscode/tasks.json` declaring an `Environment Setup` task triggered on `folderOpen` that runs `node .claude/setup.mjs`. Each file points at the other tool's directory, so each appears to belong to the other tool, and deleting one directory breaks both execution paths while leaving files behind. The payload itself sits only in `.claude/`. Opening an infected branch in either Visual Studio Code or Claude Code is sufficient to execute, with no `npm install` and no build step. ^{[10] [11] [12]}

FACT This is not a novel defect. Check Point Research disclosed three related issues in Claude Code's handling of repository-supplied project files, tracked as CVE-2025-59536 and CVE-2026-21852: arbitrary command execution through hooks declared in `.claude/settings.json` before the trust dialog was accepted, MCP server initialisation through `enableAllProjectMcpServers` before the user could read that dialog, and API key exfiltration by setting `ANTHROPIC_BASE_URL` in the same file. Disclosed 21 July 2025, fixed between 26 August and 28 December 2025, published 25 February 2026. ^[13]

ASSESSMENT The patches closed the pre-consent execution window. They did not, and could not, change the fact that project-scoped agent configuration is designed to be executable and designed to be checked into the repository. A user who accepts the trust dialog for a repository they believe is theirs has consented to whatever the worm wrote into that repository's fiftieth branch. Consent is the wrong control for content the user has not read. ^{[13] [25]}

Artefact	What it does	Which control class inspects it	Coverage
package.json dependencies	Declares packages to install	Software composition analysis	Covered
Published tarball contents	Ships the code that installs	Registry-side scanning, some SCA	Partial. ChainDrop's tarballs carried valid provenance
preinstall and postinstall scripts	Executes on npm install	SCA policy, ignore-scripts enforcement	Covered where policy is enforced
.claude/settings.json hooks	Executes a command on agent session start	None identified	Not inspected
.vscode/tasks.json folderOpen	Executes a command when the folder opens	None identified	Not inspected
.mcp.json server entries	Starts MCP servers on project load	None identified	Not inspected

ASSESSMENT The reason for the gap is a category error rather than negligence. Software composition analysis classifies a file by whether it declares dependencies. These files declare execution. They sit outside the dependency graph entirely, which is why a clean SCA report and a fully compromised checkout are compatible states. Abby Kearns of ActiveState framed the required change correctly when she told The Register to treat repository-supplied configuration as executable content, because that is what it now is. She also called ChainDrop the first campaign to notice the gap and use it at scale, and said it will not be the last. [\[12\]](#)

FACT The worm's credential harvester enumerated more than 140 hotspot paths, and the list includes .claude/credentials.json, .cursor/credentials.json, and .openai/auth.json alongside the expected .npmrc and .aws directories. It scraped GitHub Actions runner memory for secret fragments, enumerated AWS Secrets Manager and SSM across sixteen regions, and matched nineteen secret formats. Exfiltration used AES-256-GCM with a per-run key wrapped under an embedded RSA public key, and command and control resolved from an Ethereum mainnet contract across seventy-five public RPC endpoints rather than from any domain that could be seized. [\[10\]](#)

ASSESSMENT So the agent directory is simultaneously the persistence mechanism and the credential store, on a machine that by construction holds production cloud credentials and an authenticated path into the source of truth. That is the convergence worth planning for, and it does not require any advance in agent capability to become worse. It requires only that more developers keep agent configuration in more repositories, which is happening regardless.

HYPOTHESIS Offered for testing: within two quarters, repository-supplied agent configuration will be used as a primary initial access vector in a published incident, rather than as secondary persistence inside a package-manager compromise. Confirmed by an advisory describing a

repository whose only malicious content is an agent hook. Refuted by two quarters passing with the technique appearing only as a follow-on stage.

The defender advantage, treated as a hypothesis

The claim that defenders currently hold the advantage is doing a great deal of work in public discussion of this topic, and it is being attributed to Google in a stronger form than Google stated it.

FACT Writing on 13 August 2026, Daniel Fabian, who leads Google's red teams, argued that red teams should begin by automating isolated pieces of manual exercises and build modular sub-agents connected through an orchestrator. On timing he wrote that the side which operationalises agents first sets the tempo, and that the teams which lose will be the ones that treated this as a future problem while their adversaries treated it as a present opportunity. He also noted that after an initial foothold, agents often accomplish their goals before a human analyst has time to review the alert, and that concrete threat intelligence on what real-world actors are building is currently limited. The post projects open-weight models matching frontier cybersecurity capability over the next six to twelve months. ^[9]

ASSESSMENT That is a conditional advantage, contingent on acting, not a structural one. Fabian's own framing supplies the expiry date: six to twelve months from 13 August 2026 puts the open-weight parity point between February and August 2027. The advantage he describes belongs to whichever side deploys first, and the Taiwan operators had already deployed six weeks before he published. Read as a finding it is comforting. Read as written it is a deadline, and the acknowledgement that intelligence on adversary agent development is limited means nobody proposing this hypothesis can currently measure which side is ahead. ^[9]

ASSESSMENT There is a second-order problem with the argument as it is being repeated. Home-field knowledge is a real defender asset, but it is an asset in environments that are instrumented and inventoried. In an environment where nobody knew twenty-one systems were connected until an attacker mapped them, the home-field advantage accrues to whoever completes the map first, and in July that was not the defender. ^[3]

Precedent, and what it establishes

VENDOR CLAIM The Taiwan operation is not the first of its kind, and the framing of it as a first depends on the qualifier chosen. In November 2025 Anthropic disclosed a campaign it tracked as GTG-1002, in which an operator it assessed as Chinese state-linked jailbroke Claude Code and used it to automate what the company estimated at 80 to 90 percent of multi-stage intrusions against roughly thirty targets, with human involvement limited to strategic decision gates. Anthropic detected the activity in mid-September 2025 and published on 13 November 2025. ^[15]

FACT MITRE catalogued that activity as Campaign C0062 in ATT&CK, which is the closest thing this subject has to an independent, non-commercial record. The catalogue entry establishes that the campaign is treated as real by a body with no product to sell. It does not independently verify

Anthropic's 80 to 90 percent automation estimate, which remains the vendor's characterisation of its own logs. ^[16]

ASSESSMENT The useful comparison between GTG-1002 and Taiwan is not which came first. It is that in both cases the vendor holding the telemetry produced the automation percentage, and in both cases nobody else can check it. Nine months apart, two firms with different products described operations of similar shape using similar language, and the second one was received as a step change. On the published evidence it is a second data point, and two data points do not describe a trend line. ^{[15] [3]}

Scenarios to Q3 2027

Probabilities below are subjective and are the least reliable content in this report. The earliest visible signs are the useful part, because each can be monitored without access to anything private.

SCENARIO A Identity closes the door 40 percent

Enterprises and government operators fix the identity plane that every operation in the cluster relied on: phishing-resistant MFA, inventoried federation endpoints, and SSO configuration audits. Agent operations continue and stall at initial access.

Consequence. Agentic tradecraft stays a nuisance rather than a capability, and the vendor category built around agent-specific detection struggles to justify a separate line item.

Earliest visible sign. A published incident in which an agent-orchestrated operation is stopped at authentication and named as such, rather than after the fact.

SCENARIO B Volume without novelty 30 percent

The number of agent-attributed operations grows from single digits into the dozens, but the tradecraft does not change: spraying, exposed federation, unpatched middleware, and commodity OCR. Detection models hold. Analyst staffing does not.

Consequence. The binding constraint moves to triage capacity and alert volume, making automation of response, not detection, the purchase that matters.

Earliest visible sign. An incident response firm reporting agent-attributed operations as a double-digit share of engagements while reporting no new technique classes.

SCENARIO C The configuration surface goes mainstream 20 percent

Repository-supplied agent configuration becomes a routine initial access and persistence path used independently of package-manager compromise, across more than one agent runtime.

Consequence. Software composition analysis, code scanning, and developer endpoint policy all need new coverage, and the developer workstation becomes a tier-zero asset in practice rather than in theory.

Earliest visible sign. A CVE or advisory covering project-scoped configuration in a second agent runtime, plus any scanning vendor shipping agent-config inspection as a named feature.

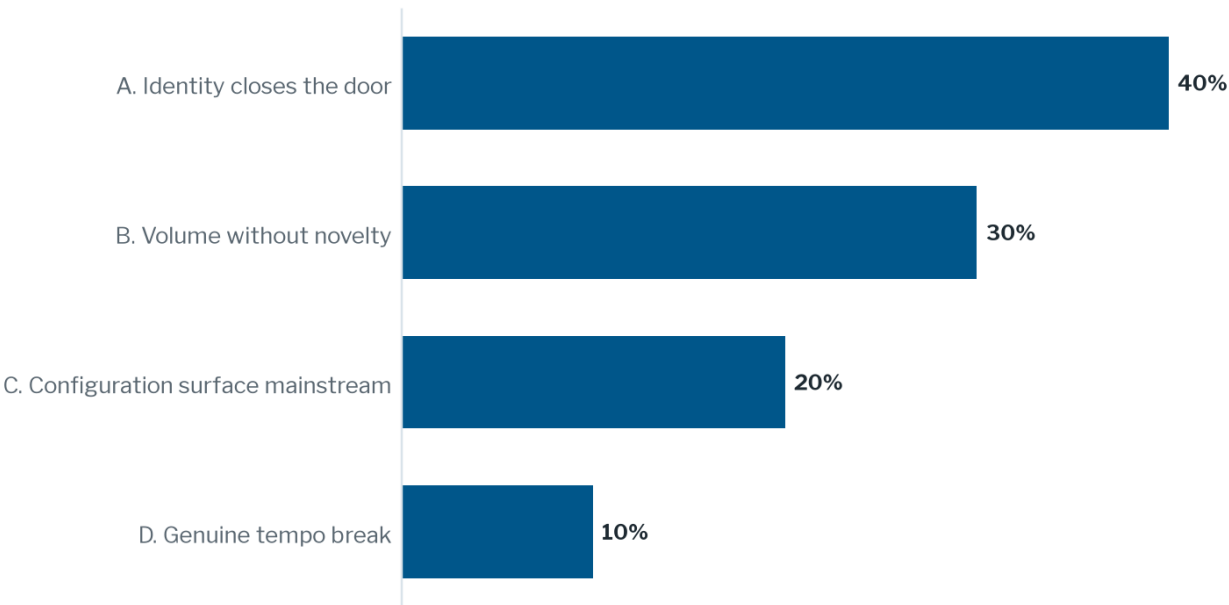
SCENARIO D A genuine tempo break 10 percent

An operation completes its objectives inside a single detection interval, minutes rather than days, against an enterprise with working instrumentation and an active SOC.

Consequence. The core argument in this report is wrong, human-paced triage becomes indefensible, and automated containment moves from optional to mandatory faster than procurement cycles can absorb.

Earliest visible sign. An incident response timeline showing initial access to exfiltration inside one hour with agent orchestration confirmed by recovered artefacts rather than inferred from behaviour.

FIGURE 7 Scenario probabilities to Q3 2027



SCENARIO Subjective probabilities from the author, not derived from a model and not the useful part of this section. They sum to 100 by construction, which is a discipline rather than a finding. Scenarios B and C are not mutually exclusive in practice.

Tripwires for the next two quarters

Each of these is observable from public sources by 31 March 2027. Together they resolve most of what this report has had to assume.

If this appears	What it means	Moves toward	Where to watch
An incident response timeline with agent orchestration and under sixty minutes from access to exfiltration	The tempo argument in this report is wrong	Scenario D	M-Trends supplements, Unit 42, CrowdStrike Global Threat Report 2027

If this appears	What it means	Moves toward	Where to watch
A scanning vendor shipping named inspection of .claude, .vscode/tasks.json, or .mcp.json	The coverage gap is being closed commercially	Scenario C, then defused	Release notes from Snyk, Socket, Semgrep, GitHub Advanced Security
A worm or intrusion set using agent configuration as the primary vector, not secondary persistence	The surface has become first-class initial access	Scenario C	npm and PyPI advisories, GitHub Security Lab
An open-weight model matching frontier cyber capability on a public benchmark	Google's six-to-twelve-month projection resolves early	Scenario D	Model release notes, Cybench and CTF benchmark results
A CISA or ENISA advisory on agentic intrusion sets carrying detection logic, not awareness language	Vendor claims have been converted into public doctrine	Scenario A	CISA advisories, ENISA threat landscape publications
A government publishing forensic artefacts rather than a summary statement	The evidence base gets a non-commercial floor	Resolves the central limitation here	Taiwan MODA and NICS, national CERT bulletins
An insurer filing a named agentic-attack sublimit or AI-attack exclusion	The exposure is being priced as distinct	Scenario B or D	Lloyd's market bulletins, admitted-market form filings
No new named agentic operations by 31 March 2027	The cluster was a reporting artefact of one quarter	Falsifies the premise of this report	Tenable RSO, Unit 42, Mandiant

Implications

CISOs, security architects, and detection engineering leads

ASSESSMENT Do not buy an agent-speed detection product this budget cycle. Nothing in the published record shows an agent operation moving faster than the human tradecraft your controls are already calibrated against, and the one government-confirmed timeline in evidence shows a sixteen-day detection lag on a four-day operation, which is a coverage failure that no amount of detection speed corrects. Spend instead on the thing that failed: authentication telemetry from every federated system, actually queried. Tenable, Mandiant, and the Taiwan case independently point at the same entry class, and it is identity, not novelty. Tenable's own defensive guidance, discounted for the fact that it sells the platform, lands on phishing-resistant MFA and least privilege rather than on anything agent-specific. [\[1\]](#) [\[5\]](#) [\[6\]](#) [\[8\]](#)

ASSESSMENT For detection engineering specifically, the change to make is in what gets modelled rather than how fast. Agent operations are volumetrically regular and abnormally patient, which are properties you can write rules against: authentication attempt cadence with

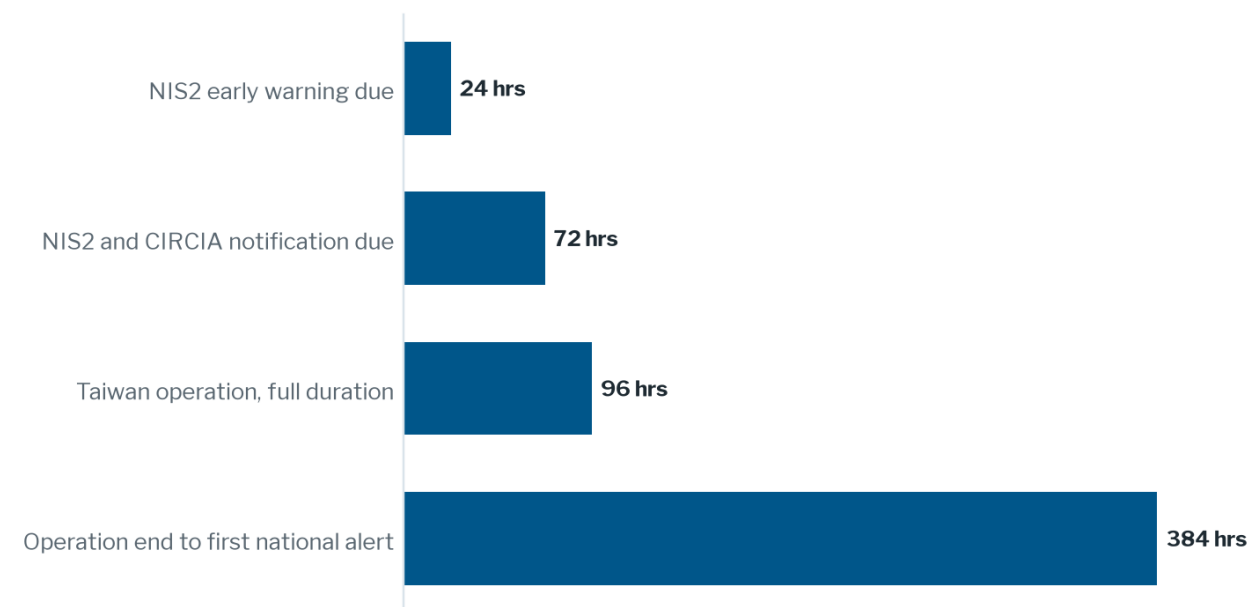
machine-like inter-arrival regularity, breadth of distinct systems touched by one identity within a session, and enumeration depth against a single API surface. Those signatures suit an operator working eight sub-agents and do not suit a human. Whether they survive an adversary who deliberately introduces jitter is an open question, and one worth testing on your own red team before an outsider tests it for you. ^[3]

ASSESSMENT Add the developer estate to the detection scope. A checked-in `.claude/settings.json` or `.vscode/tasks.json` that nobody on the team authored is currently invisible to the entire security stack, and it executes on folder-open. File integrity monitoring for those paths on developer endpoints costs nothing and closes the specific gap ChainDrop used at scale. ^{[10] [12]}

Critical infrastructure operators and their regulators

ASSESSMENT The uncomfortable arithmetic for this audience is that the Taiwan operation ran to completion in ninety-six hours. Under NIS2 an early warning is due within twenty-four hours and a full notification within seventy-two. Under CIRCIA, once final, a substantial incident report is due within seventy-two hours. The operators finished, exfiltrated, and withdrew inside a window in which a covered entity that had detected the intrusion on day one would still have been drafting its first regulatory notification. Taiwan did not detect it on day one. It detected it on day nineteen. ^{[1] [23] [24]}

FIGURE 8 Reporting windows against the observed timeline



ASSESSMENT Regulatory deadlines are as summarised in secondary legal commentary, not read from the instruments themselves, and CIRCIA's final rule was still pending at the time of writing. The operational hours come from the Ministry of Digital Affairs timeline. The comparison is illustrative: reporting clocks start at detection, not at intrusion. ^{[1] [23] [24]}

ASSESSMENT The supply chain dimension matters more here than the agent dimension. The operation pivoted from an initial government department to the nuclear safety agency, to government IT supply chain vendors, and to at least seven energy companies. Whatever the tooling, the pivot worked because those relationships were reachable and mapped. An operator

inventory of which vendors hold federated access into your OT-adjacent IT estate is worth more than any agent-specific control, and most operators reading this cannot produce one within a day. ^{[3] [4]}

ASSESSMENT For regulators, the specific request the evidence supports is publication of forensic artefacts rather than confirmation statements. Taiwan's ministry has given the only non-commercial account of this incident, and it consists of four sentences. Every operational detail the sector is now planning against comes from a firm that sells preemptive defence. A national CERT publishing indicators, timelines, and telemetry would do more to improve sector defence than a further year of vendor threat reports, and it is within existing authority. ^[1]

Developer platform, DevSecOps, and supply chain owners

ASSESSMENT You have two separate problems from ChainDrop and only one of them is being discussed. The first is that provenance attested malware correctly, which means SLSA and trusted publishing do not answer the question your board thinks they answer. Adopt minimum release age as an actual gate. The npm CLI supports it from version 11.10.0, and a three-to-seven-day cooldown would have contained a worm that did its damage in four hours and eighteen minutes. ^{[10] [11] [14]}

ASSESSMENT The second problem is the one nobody has a control for. Treat every project-scoped agent configuration file as executable content under change control: `.claude/settings.json`, `.mcp.json`, `.vscode/tasks.json`, and their equivalents in whatever agent runtimes your developers have installed without telling you. Inventory them, pin them, require review on change, and scan all branches rather than the default one, because ChainDrop planted across as many as fifty. Your SCA tool will report clean throughout, and that report will be accurate and useless. ^{[10] [11] [12]}

ASSESSMENT Reconcile the registry tarball against the git tree for your critical dependencies. ChainDrop's entire propagation model depended on those two diverging without anyone comparing them. This is not a product category yet, which means it is a scripting job that a platform team can do in a sprint and a control that most competitors will not have for a year. ^[11]

Cyber insurance underwriters and risk quantification teams

ASSESSMENT The exposure that has actually changed is correlation, not severity. A four-day intrusion against one agency is an ordinary loss. A worm that reaches 444 packages across a dozen victim organisations in under four and a half hours, planting persistence in developer environments that no scanner inspects, is a correlated event across an entire book, and the current generation of aggregation models does not have a parameter for repository-supplied execution. Price the developer estate as an accumulation zone. ^{[10] [11]}

THIRD-PARTY ESTIMATE Market context, all of it soft. Cyber premium is projected around 16.4 billion dollars globally for 2026, rates that softened for two years are expected to climb 15 to 20 percent, and the World Economic Forum's 2026 outlook names AI-driven threats as a leading systemic risk. Some carriers introduced AI security riders in 2026 requiring proof of red-teaming

and documented risk assessment. These figures come from broker and reinsurer commentary with an obvious interest in hardening, and should be treated as directional only. ^{[20] [21] [22]}

ASSESSMENT The underwriting question worth adding to the application is not whether the insured uses AI. It is whether the insured can enumerate, on request, every project-scoped agent configuration file in its repositories and say who authored each one. An insured that can answer has a control the market has not priced yet. An insured that cannot has an unbounded execution surface on machines holding production credentials, and no scanner reporting on it.

National security policy staff and standards bodies

ASSESSMENT The policy risk in this episode is not the capability. It is that a sovereign incident is being characterised for the international record by a commercial vendor. Taiwan confirmed the intrusion in four sentences and published nothing operational. Dream published 1,395 files worth of interpretation. Every subsequent policy discussion, including this report, is built on the vendor's reconstruction because no state alternative exists. That is a durable structural problem and it will recur with every incident of this kind. ^{[1] [3] [4]}

ASSESSMENT Attribution discipline held reasonably well here and deserves noting. Dream said Chinese-language operator materials and stopped. Tenable assessed state-adjacent contractor or patriotic hacker origin as the leading explanation and said so with hedging intact. Taiwan's ministry did not name China at all. The gap opened downstream, in headlines that converted a language artefact into a state actor. Standards bodies can close some of that gap by defining a shared vocabulary for autonomy levels in intrusions, because first fully autonomous attack currently means whatever the writer wants it to mean, and the government that was attacked described the same operation as a hybrid of manual and agent-assisted work. ^{[1] [3] [5]}

ASSESSMENT The concrete standards gap is a classification scheme for agent autonomy in offensive operations, with observable criteria: whether human decision gates existed, how many, at what points, and whether the framework selected targets or executed against a supplied list. Without it, every incident of this kind will be reported at the highest autonomy level the evidence can be stretched to support, and defenders will keep procuring against a threat model set by press releases.

What this document cannot answer

Six questions matter to the argument above and cannot be answered from public sources.

Whether the Dream archive is what it appears to be. Nobody outside the firm has examined the 1,395 files. Independent examination, or publication of hashes and a subset of artefacts, would settle the question of authenticity and the more interesting question of how much human direction the framework actually needed between waves.

How many human decision gates the Taiwan operation contained. Taiwan's ministry described a hybrid of manual and agent-assisted operation. Dream described autonomous waves with self-

correction. Both can be true at different granularity, and the difference between five decision gates and fifty is the difference between an autonomy story and a productivity story.

What the four unnamed incidents in Tenable's cluster are. Without them the count cannot be checked and the implied rate of increase cannot be evaluated.

Whether any software composition analysis or code scanning product currently parses project-scoped agent configuration for execution content. No documentation reviewed for this report says so and no vendor was asked directly. A short round of vendor questioning would resolve it and is the single highest-value piece of primary research available on this subject.

How many repositories currently carry agent configuration files authored by someone other than the repository owner. ChainDrop planted across as many as fifty branches per repository and infected packages were yanked, but the config files sit in git history and in every clone taken during the window. A survey of public repositories would produce a real number and nobody has published one.

Whether detection tooling can distinguish agent-driven authentication behaviour from human-driven behaviour at useful precision. The signatures proposed in this report are plausible and untested. Only an organisation with a red team running agent frameworks against its own instrumented estate can answer it, and those that have are not publishing.

Half-life of this assessment

Decays within six months, by February 2027. Every operational count attributed to Dream, because independent examination or a competing reconstruction will either firm them up or discredit them. The seven-incident cluster figure, which will be superseded by a larger number whose composition is equally unverifiable. The claim that no scanner inspects agent configuration files, which is the kind of gap a vendor closes in one release cycle once it is named in the trade press. All five survey percentages, which are annual instruments with new fieldings due.

Decays within twelve months, by August 2027. The comparison of agent tempo against human breakout time, which depends on CrowdStrike's 2026 figures and on agent frameworks not improving, and the second of those is a poor bet. Google's projection of open-weight parity in frontier cyber capability resolves inside this window by its own terms. The scenario probabilities should be regarded as void after Q3 2027 regardless of what has happened.

Decays within twenty-four months. The specific tooling names: Hermes, OpenClaw, Tesseract, the current agent runtimes and their configuration file paths. Attack frameworks and agent products churn faster than this, and a reader in 2028 should treat every product name in this report as a period detail.

Architectural, and unlikely to decay. That provenance attests the path from source to artefact and cannot attest authorship, which is a property of the design rather than of any implementation. That repository-supplied configuration which executes is executable content regardless of what

the file extension suggests, and that any control classifying files by whether they declare dependencies will miss it. That detection latency measured in days makes an adversary's tempo irrelevant below the resolution of that latency. And that when the only detailed account of a sovereign incident comes from a firm selling the remedy, the resulting threat model is shaped by a commercial interest, whether or not any individual fact in it is wrong.

Limitations

This report was written independently. No vendor named here commissioned, reviewed, funded, or was given advance sight of it, and the author holds no position in any company mentioned and has no commercial relationship with any of them.

The evidence base is weak in a specific and important way. Of the material facts underpinning the analysis, four are checkable against non-commercial primary sources: Taiwan's ministerial statement, the two CVE records, MITRE's catalogue entry for the 2025 campaign, and the published documentation for the agent configuration formats. Everything else, including every operational count, every telemetry median, and every survey percentage, comes from firms that sell products addressing the problem those numbers describe. This report labels that distinction rather than resolving it, because it cannot be resolved from outside.

The author did not examine the Dream archive, did not receive briefings from Dream, Tenable, StepSecurity, Unit 42, or Check Point, and did not put questions to any scanning vendor about coverage of agent configuration files. That last omission is the most consequential, and it is named in the section on what this document cannot answer for that reason.

Two comparisons drawn here are structurally imperfect and should not be over-read. Setting a single incident's duration against population dwell-time medians compares a case to a distribution. Setting eCrime breakout time against total operation length compares two differently defined intervals, one vendor-specific and one derived from a timeline. Both were used because they are the best available and both would be improved by a like-for-like measurement that nobody has published.

Regulatory deadlines cited for NIS2 and CIRCIA were taken from secondary legal commentary rather than from the instruments, and CIRCIA's final rule had not been issued when this was written. They are used to establish orders of magnitude, not compliance positions, and no reader should take a filing decision from them.

Finally, the second-order argument about agent configuration files is the least corroborated part of this report and the part the author is most confident about, which is an uncomfortable combination and is stated here so a reader can discount it accordingly.

Sources

- [1] Taipei Times. "AI-driven hacking campaign targets Taiwan government agencies." 13 August 2026. Reporting the Ministry of Digital Affairs statement. *Independent press* [taipetimes.com](https://www.taipetimes.com)

- [2] NBC News. "Taiwan says it was targeted last month in AI-driven hacking campaign." 13 August 2026. *Independent press* [nbcnews.com](https://www.nbcnews.com)
- [3] The Register. "'Near-autonomous' AI agents attack Taiwan's nuclear safety agency." 12 August 2026. *Independent press* [theregister.com](https://www.theregister.com)
- [4] CyberScoop. "Researchers observe first 'near-autonomous' AI attack on government target in Taiwan." August 2026. *Independent press* [cyberscoop.com](https://www.cyberscoop.com)
- [5] Tenable Research Special Operations. "The Agentic AI Threat Cluster: Seven Incidents, Three Actors, and What They Mean." 2026. *Vendor primary (commercial interest)* tenable.com
- [6] Tenable. "Agentic AI Attacks: A Practical Defense with Exposure Management." 2026. *Vendor primary (commercial interest)* tenable.com
- [7] CrowdStrike. "2026 Global Threat Report." February 2026. *Vendor primary* [crowdstrike.com](https://www.crowdstrike.com)
- [8] Mandiant, Google Cloud. "M-Trends 2026: Data, Insights, and Strategies From the Frontlines." March 2026. *Vendor primary* cloud.google.com
- [9] Fabian, Daniel. "The Evolving Role of the Red Team in the Era of Agentic Security." Google Security Blog, 13 August 2026. *Vendor primary* blog.google
- [10] StepSecurity. "ChainDrop npm Worm: Bun-loaded CI/CD credential harvester with Ethereum dead-drop C2." August 2026. *Vendor primary (commercial interest)* stepsecurity.io
- [11] Unit 42, Palo Alto Networks. "ChainDrop: Inside a Self-Propagating npm Worm." August 2026. *Vendor primary (commercial interest)* unit42.paloaltonetworks.com
- [12] The Register. "ChainDrop worm crawls into npm supply chain, evades standard defenses." 15 August 2026. *Independent press* [theregister.com](https://www.theregister.com)
- [13] Check Point Research. "Caught in the Hook: RCE and API Token Exfiltration Through Claude Code Project Files. CVE-2025-59536, CVE-2026-21852." 25 February 2026. *Vendor primary* research.checkpoint.com
- [14] npm. "Trusted publishing for npm packages." npm Docs. *Vendor primary* docs.npmjs.com
- [15] Anthropic. "Disrupting the first reported AI-orchestrated cyber espionage campaign." 13 November 2025. *Vendor primary* anthropic.com
- [16] MITRE. "Anthropic AI-orchestrated Campaign, Campaign C0062." MITRE ATT&CK. *Standards body primary* attack.mitre.org
- [17] Kai and Wakefield Research. "2026 State of Autonomous Defense Report." Fielded 15 to 29 June 2026, n=500 CISOs at companies above 500 million dollars revenue. *Vendor-sponsored survey* prnewswire.com
- [18] Futurum Group. "Is the Rise of Agentic AI Threatening Cybersecurity Readiness?" Citing 1H 2026 Cybersecurity Global Enterprise Decision Maker Survey, n=929. *Analyst firm* futurumgroup.com
- [19] Sygnia. "CISO Survey 2026: The State of Incident Response Readiness." April 2026. *Vendor-sponsored survey* sygnia.co
- [20] World Economic Forum. "Global Cybersecurity Outlook 2026." January 2026. *Standards body primary* weforum.org
- [21] Munich Re. "Cyber insurance: Risks and trends 2026." *Vendor primary* munichre.com
- [22] Insurance Business. "Cyber insurance enters the AI risk era as limits, wording and underwriting models shift." 2026. *Independent press* insurancebusinessmag.com
- [23] Consilium Law. "The 72-Hour Clock Is Coming: CIRCIA's Mandatory Cyber Incident Reporting." 2026. Secondary summary; the final rule had not been issued at the time of writing. *Second-hand (unverified)* consilium.law
- [24] Legiscopes. "NIS2 Incident Reporting: The 24h/72h Framework." Secondary summary of Directive (EU) 2022/2555 reporting obligations. *Second-hand (unverified)* legiscopes.com
- [25] Anthropic. "Claude Code settings." Claude Code documentation. *Vendor primary* code.claude.com

Rights and permitted use

© 2026 David Soden, davidsoden.com/market-assessment. All rights reserved.

This document, including its structure, analysis, findings, evidence classification, and framing, is the original work of David Soden, published at davidsoden.com/market-assessment.

It is published for public reading. You are free to share the complete, unmodified file, to link to it, and to quote short passages in commentary, reporting, or discussion, provided the author and the site above are credited.

The following require prior written consent: republishing this work in whole or in substantial part under another name, masthead, or brand; excerpting or modifying it in a way that alters the analysis or removes the evidence classifications; incorporating any portion of it into a paid, commissioned, or client-facing deliverable; resale or distribution behind a paywall; and use of this document as training data for machine learning models.

Commissioned Market Assessment reports, commercial licensing, and briefings on this material are available.

This document is analysis, not investment, legal, or procurement advice. It was not commissioned, reviewed, or funded by any company named in it, and the author holds no position in any of them.

Market Assessment reports, commissioned research, and briefings: davidsoden.com/market-assessment