

The background features an abstract geometric pattern of yellow lines and dots on a dark blue background. In the top right, a star-like shape is formed by several lines meeting at a central point. In the bottom left, a complex pattern of concentric circles and intersecting lines surrounds a central yellow circle with a cross inside. Other scattered dots and lines create a sense of depth and connectivity.

Claude runs Reco. Reco secures Claude.

Complete AI security, both directions.

Complete visibility across every Claude surface and across everything Claude connects to.

Now listed by Anthropic. Reco's integration with the Claude Compliance API is published in Anthropic's integrations listing.

The security gap

Claude is already in your enterprise. Your security stack probably doesn't see it yet.

Employees use it daily through Claude Enterprise. Developers build on it through the Anthropic Console. Anyone in the organization can now deploy autonomous agents that take actions, access data, and connect to your other systems without human sign-off on every step.

Most security tools cover one of these surfaces. Reco covers all three. And unlike tools that stop at the Claude boundary, Reco correlates what happens inside Claude with what happens across every SaaS app it touches.

Six risk categories driving Claude security conversations

1 Shadow AI and unauthorized usage

Claude is adopted by individual teams without central approval. Shadow AI is now the third most common non-malicious insider action in DLP datasets, up fourfold year-over-year.

2 Sensitive data leakage

Enterprise data flows into AI platforms at scale. The risk isn't just leakage; it's systematic extraction through AI-assisted workflows that existing controls weren't built to catch.

3 Prompt injection

Threat actors deliver malicious instructions to agents via chat, malware, or embedded content. Prompt injection is on the OWASP LLM Top 10 with real-world public incident write-ups.

4 API key exposure

Long-lived API keys embedded in scripts persist beyond their original use. A stolen Anthropic API key gives attackers access to shared project files, unauthorized generation, and cost abuse.

5 Excessive agent agency

Autonomous agents are live in production: orchestrating workflows, accessing sensitive data, making decisions traditional controls were never designed to catch.

6 Access control and offboarding gaps

If Claude access isn't tied to role-based policies, former employees and overpermissioned users retain access long after they should.

Two surfaces for every agent. Most tools cover one.

Claude runs across two administrative surfaces. Many security tools cover only Claude Enterprise via the Compliance API. Reco covers both:

Claude Platform (the developer console)

Where engineering teams manage API keys, workspaces, and workspace membership. This is the developer surface, and it's where API key exposure and unauthorized access typically begin. It's also where Managed Agents run: the agents engineering teams build and deploy on Anthropic's infrastructure. Reco maps each agent's model, version history, tools, permission policies, and every MCP server it connects to.

Claude Enterprise

Where employees use Claude day-to-day, governed via the Compliance API: user directory, groups, roles, permissions, projects, and the activity feed.



That multi-surface coverage means a CISO can govern Claude the same way they govern Okta, Salesforce, and Microsoft 365 – same identity graph, same posture checks, same incident response workflow.

What makes Reco different

While most tools see Claude in isolation, Reco sees Claude in context which is where the real risk lives.



Cross-source identity graph

Reco correlates Claude activity with signals from 230+ connected SaaS apps, your IdP, endpoints, and network – building a unified picture of what every user and agent is doing across your entire estate, not just inside Claude.



Toxic combination detection

Reco surfaces dangerous combinations that no single-source tool can see: an overpermissioned agent whose owner left the company, connected to a sensitive SaaS app, with no MFA. Each signal alone looks benign. Together, they're a critical risk.

Example: an agent with access to your Salesforce CRM and Google Drive looks unremarkable on its own. But when Reco Graph shows the agent's owner left the company last month, its permissions were never revoked, and it has no guardrails set – that's a toxic combination. Reco surfaces it. Other tools don't.

How it works

Reco governs Claude

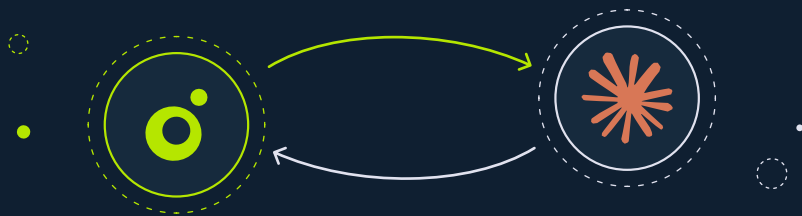
Reco connects to all of Claude’s surfaces and gives your security team continuous posture management, identity governance, and threat detection across all of them.

When something changes – a new API key, an agent granted new permissions, a former employee still active in Claude Enterprise – Reco surfaces it immediately. Your team knows who built each agent, whether that person still works there, what the agent can reach, and what changes the moment something goes wrong.

Claude queries Reco

Reco exposes its own MCP server. That means Claude connects to Reco as a data source and queries it directly. Your security team opens Claude, asks a question in natural language, and gets the answer from Reco’s normalized identity graph.

The queries your team runs through Reco’s MCP are real security investigations: who has access to what, what changed, what’s anomalous.



Full capability coverage

Each capability maps to a business outcome and something you can demonstrate to auditors and your board.

CAPABILITY	WHAT THIS MEANS FOR YOUR BUSINESS	YOU CAN PROVE TO AUDITORS
Shadow AI discovery	Know every Claude account in your org sanctioned or not before auditors do. Seven sensors surface usage that authenticated-users-only tools never see: endpoints via EDR (CrowdStrike, Defender, SentinelOne), network via SASE, browser via browser guard, plus SSO, social logins, plugins, and shadow email. Every app classified as sanctioned, shadow, GenAI, or AI agent.	Who is using Claude, on what device, and connected to what data.
Posture & configuration	Catch misconfigurations across all of Claude’s surfaces before they become incidents. 100+ one-click checks, drift detection, and custom policies via Policy Studio – mapped to SOC 2, ISO 27001, CIS, NIST, PCI DSS, and HITRUST.	Your Claude deployment meets the controls your auditors require.

CAPABILITY	WHAT THIS MEANS FOR YOUR BUSINESS	YOU CAN PROVE TO AUDITORS
AI governance	Govern Claude sessions, projects, and GenAI usage across your org. Map every agent back to its owner and the data it can reach. Surface shadow accounts, enforce MFA, and flag stale or overpermissioned access – with IdP correlation across Okta, Azure AD, and Ping.	Only the right people have access and access ends when employment does.
Agent security	Red team and assess every agent built on Claude prompt analysis, access, tools, activity, permissions, guardrails, and blast radius. Enforces least privilege for non-human identities across Reco's cross-source identity graph.	No agent has more access than its job requires.
Threat detection & response	Detect account compromise, data exfiltration, and risky agent behavior across Claude and the 230+ apps it connects to. 400+ pre-built detections and behavior analytics, correlated into incident narratives, with automated response through your existing SIEM and SOAR.	Threats are caught and responded to, not just logged.

Why Reco

Most tools stop at the Claude Enterprise boundary and treat Claude as an isolated product to monitor. Reco treats Claude as part of your broader identity and access estate – because that's how attackers treat it.

With 230+ integrations already in place, Reco correlates Claude activity with what's happening in Salesforce, Google Workspace, GitHub, Slack, and every other app your agents touch. That cross-source context is what makes toxic combination detection possible – and what makes Reco different from tools that only see inside Claude.

Close the security gap in your expanding AI footprint.

Request a demo at reco.ai

Request a Demo