

Monitor, Educate, and Enforce AI Policy at the Browser

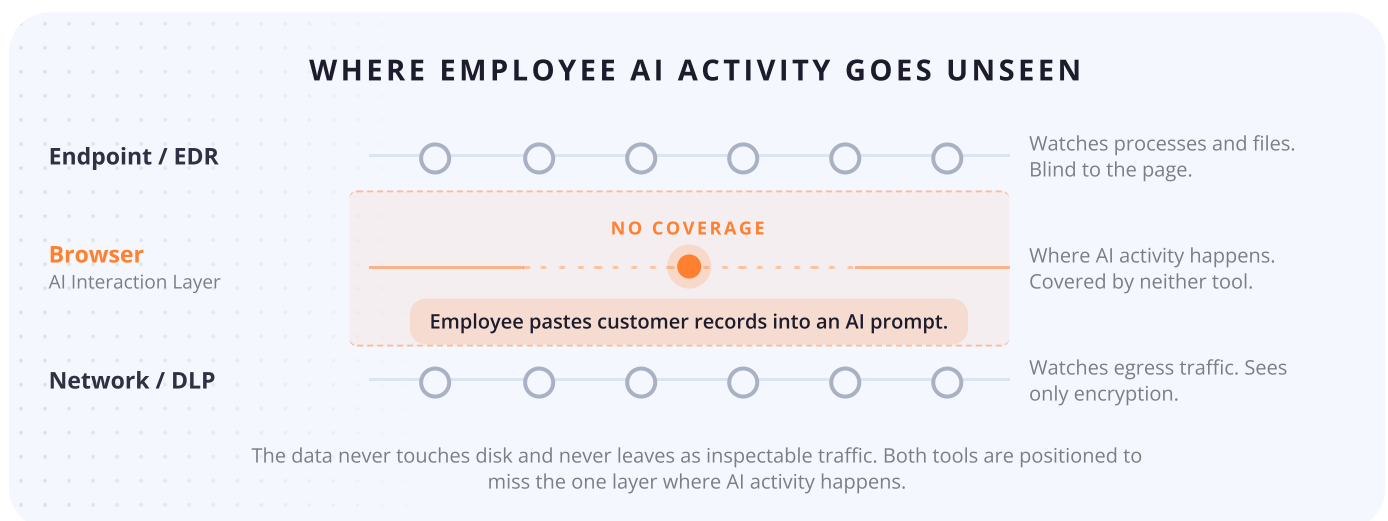
How Browser-Level Visibility Covers the AI Activity That Endpoint and Network Tools Alone Cannot See

THE REALITY

Employees access AI tools through their browser, drafting documents, analyzing data, and writing code with minimal governance coverage. Indeed, [industry research](#) found that 80% of employees use unapproved AI tools while only 36% of organizations maintain governance policies covering that use. The resulting exposure is measurable. A [recent workforce study](#) discovered that 77% of employees have pasted company data into AI tools, most of it through personal accounts invisible to the security team.

Predictably, written policy has proven insufficient on its own. [HR Dive reported](#) that 45% of surveyed employees have used AI tools their companies explicitly banned, and 58% have pasted sensitive data into them. The exposure this creates is well documented. In one of the most cited incidents, Samsung engineers [leaked confidential source code to ChatGPT three times in under a month](#), prompting a company-wide ban. The ban solved neither problem, since usage continued through personal devices while the visibility that might have caught it disappeared entirely.

THE BLIND SPOT



The limitation is architectural. Endpoint detection and response platforms monitor processes, file systems, memory, and system calls, but an employee pasting customer records into an AI prompt field triggers none of them. The interaction happens entirely inside the rendered page, leaving no trace at the file system or process level.

[Analysis of EDR platforms](#) confirms that leading endpoint solutions cannot access browser information, leaving them unable to observe the layer where AI interaction concentrates. [Browser security research](#) puts the same conclusion plainly, finding that endpoint and network security tools are failing to keep pace with threats operating inside the browser.

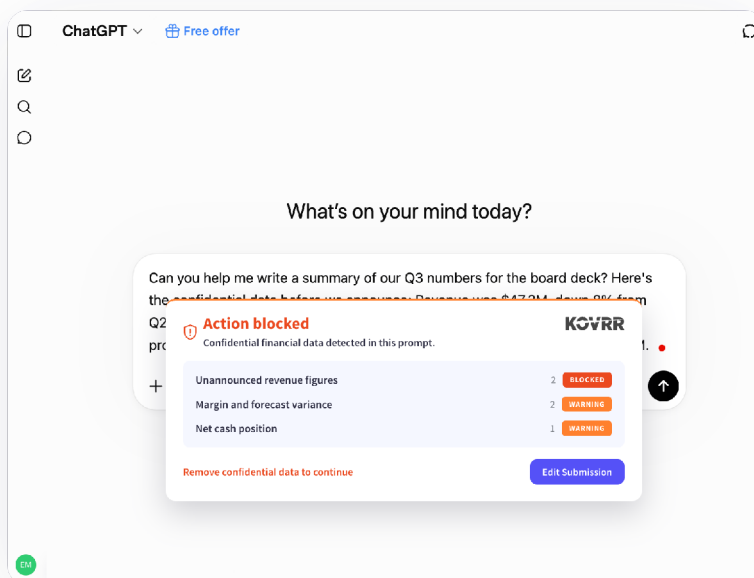
Data loss prevention tooling carries the same structural constraint. Legacy DLP relies on [file scanning and pattern matching at network egress points](#), inspection stages that a copy-paste action into a prompt field never crosses. The data leaves through an encrypted session initiated by a sanctioned browser, indistinguishable at the network layer from ordinary web activity.

Meanwhile, the volume moving through that channel keeps growing. [Enterprise AI usage research](#) found that 39.7% of all AI interactions involve sensitive data, with the average employee submitting proprietary information to an AI tool once every three days. [Parallel analysis](#) covered by The Hacker News determined sensitive data in more than 6% of enterprise AI conversations, with some platforms exceeding 12%.

Blocking has emerged as the reflexive response, and it consistently backfires. A blocked domain pushes usage onto personal accounts and personal devices, where corporate security has no reach at all. [More than half of employees admit](#) to using unapproved AI tools at work, with one in ten knowingly sharing sensitive data through them.

In the end, the organization loses the telemetry while keeping the exposure. Effective policy enforcement requires presence at the point of interaction, with the judgment to distinguish an acceptable prompt from a dangerous one in real time, and traditional security architecture has no component positioned to do it.

THE PLATFORM RESPONSE: AI INTERACTION DATA FABRIC



The limitation is architectural. Endpoint detection and response platforms monitor processes, file systems, memory, and system calls, but an employee pasting customer records into an AI prompt field triggers none of them. The interaction happens entirely inside the rendered page, leaving no trace at the file system or process level.

[Analysis of EDR platforms](#) confirms that leading endpoint solutions cannot access browser information, leaving them unable to observe the layer where AI interaction concentrates. [Browser security research](#) puts the same conclusion plainly, finding that endpoint and network security tools are failing to keep pace with threats operating inside the browser.

[Kovrr's AI Security and Governance Platform](#) operates precisely where the interaction happens. The platform surfaces every AI interaction at the point of use through browser-level telemetry, identifying which applications employees engage with, which accounts they use, and which categories of data flow into each prompt. Detection runs against a catalog of more than ten thousand AI applications and over five hundred validated data categories spanning every major regulation and region, using deterministic pattern matching rather than inference.

Pre-built detection profiles group employees by department, role, and location, so security teams can apply differentiated policy from the moment the platform deploys rather than building rules for each individual by hand. Deployment runs through standard MDM, with no proxy and no browser replacement required.

Enforcement then happens in the moment. When a prompt contains PII, credentials, or regulated data, adaptive policy responds before submission, educating employees on why a prompt violates policy and guiding them toward compliant alternatives, or stopping the action outright depending on the configured rule.

Every interaction becomes a training moment, helping employees move toward safer AI habits over time and turning each near-miss into behavior change. Employees keep working inside the tools they have chosen, and policy stops residing merely in a document nobody reads. The architecture also respects privacy by design, sending only categorical signals to the platform while prompt content stays on the endpoint.

Events captured at the browser feed the broader Kovrr platform in real time. AI applications discovered through browser telemetry populate the [AI asset inventory](#) automatically. Usage patterns inform risk scoring in the [AI Vendor Risk Catalog](#), and interaction-level data flows into AIRQ ([AI Risk Quantification](#)), where browser-layer exposure translates into financial terms. Policy enforcement, discovery, and quantification operate on the same connected telemetry rather than in separate consoles.

THE RESULT

With Kovrr's platform, security teams gain enforcement at the layer where employee AI activity concentrates and where existing tooling has no presence. AI policy becomes an operational control instead of a written aspiration, applied consistently across every AI application at the moment data would otherwise leave the organization. Employees retain the productivity that drove them to AI tools in the first place, and every enforcement event doubles as a training moment that moves the workforce toward safer AI habits over time. Discovered applications and captured interactions strengthen the rest of the platform, keeping the AI inventory current and the financial exposure picture accurate.

Bring every AI interaction into view, and ensure policy is enforced before data leaves. [Book a demo.](#)

KOVRR

AI INTERACTION DATA FABRIC

*Telemetry alone is data. Telemetry triangulated across sources
is evidence.*