

Semantic DLP

The Smarter Way to Secure AI

GenAI has been transforming the business world since its debut just a few years ago. Innovation has taken off, and productivity is accelerating. The dreaded role of meeting notetaker? Gone. That end-of-day proposal? Finished before your coffee gets cold. Seriously, what's not to love?

You know AI can deliver real value to your organization, and you want your teams using it. But you also know your data security governance isn't really keeping pace. When sensitive data gets shared with AI tools, the risk multiplies fast. After all, AI doesn't just use your data — it amplifies it.

The challenge gets even trickier as new data is created every day, sometimes by these same AI tools. Fresh records containing sensitive information often go unlabeled, leaving permissions and access controls behind. And once that data is absorbed into an AI model, control becomes a lot murkier. Where does it go? Who might see it next?

Locking It Down Pushes Risk Underground

The instinct to clamp down is understandable. But too much governance creates a different kind of risk. When control comes at the expense of usability, problems don't disappear — they just go underground. Employees won't stop using AI; they'll simply stop telling you about it.

Effective governance requires a more balanced approach that meets employees where they are, protects what matters most, and adapts as usage and risk evolve.

Key Benefits

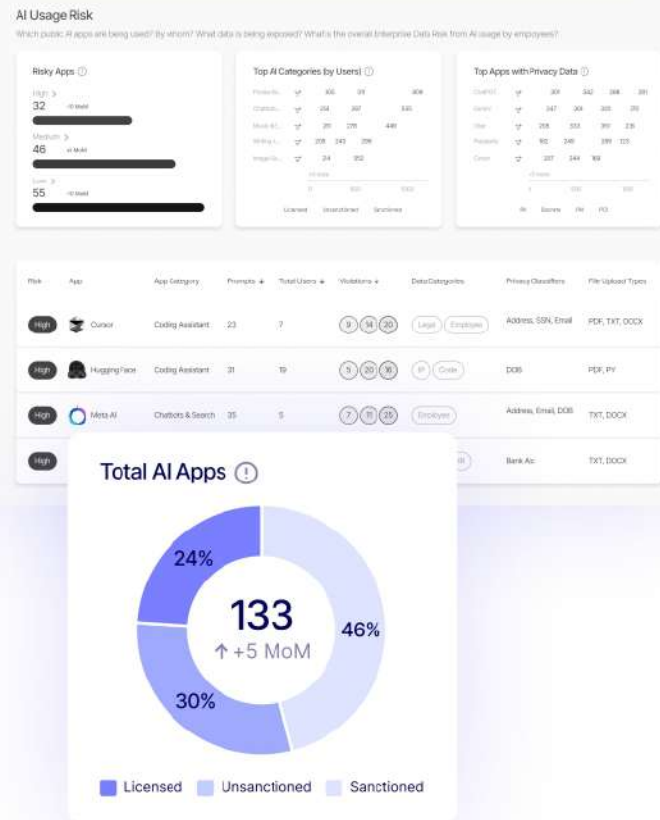
- Uncover shadow AI
- Establish guardrails for safe AI use, leveraging contextual discovery and labeling from Semantic Intelligence™
- Prevent sensitive data from being shared with AI tools
- Autonomously classify and protect new data as it appears
- Identify your riskiest applications and users
- Track all prompts, responses, and policy violations in granular detail

Shining a Light on Shadow AI

So, you've written an acceptable use policy outlining which AI applications employees should use. But how confident are you that everyone's actually following it?

Semantic DLP gives you clear visibility into all the public AI tools users are accessing, not just the usual suspects like ChatGPT, Perplexity, and Claude, but also the newer, niche tools that seem to appear overnight and are often introduced by interns, new grads, or even someone's tech-savvy teenager.

And visibility is just the beginning. Semantic DLP helps you quickly identify unsanctioned applications, aka shadow AI, and see exactly how they're being used: how many users are involved, how often the tools are accessed, how many policy violations are occurring, and which types of sensitive data are being shared. Armed with these insights, you can take action by setting smart controls or blocking risky tools entirely at the network or endpoint level.



When Context Powers Control

Now that you've got visibility into what's happening across your environment, the next step is to set guardrails to keep sensitive data where it belongs — and out of AI prompts. As with any technology that relies on data labels, the effectiveness of your AI data security ultimately comes down to how accurate those labels are.

With Concentric AI Semantic Intelligence in place, you're ahead of the curve. Sensitive data across your on-prem and cloud environments has already been discovered, categorized, and labeled using proprietary AI that understands not just what the data says, but what it means.

The result? AI controls and policies built on a solid, trustworthy foundation, so enforcement is precise, effective, and reliable.

And because Semantic Intelligence continuously monitors data, any new sensitive content generated by AI is automatically identified, labeled, and protected without manual intervention.



Guardrails Your AI Can't Ignore

With Semantic DLP, you're in control of which data stays out of public AI tools. Set up policies that restrict specific data types from being shared, whether you want one rule for everyone or controls tailored by department or application. And when someone tries to sneak past the rules, you decide the outcome: show a polite warning, redact sensitive content on the fly, or block the action entirely.

Even better, the pop-up alerts that appear when a policy is triggered can be customized to your branding and messaging, complete with links to acceptable use policies. That way, your users aren't just stopped; they're educated and nudged toward smarter choices in real time.

AI Usage Guardrails

Define guardrails for safe and compliant use of AI services by your enterprise users, leverage guardrails from the out-of-the-box library or create your own.

Search...

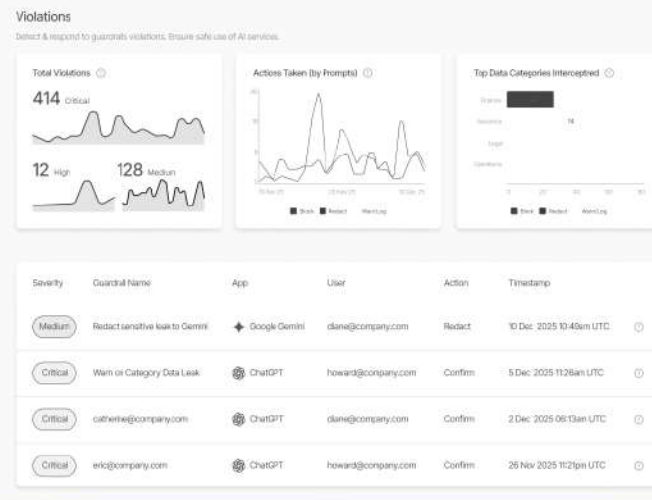
Notification Settings

Severity	Guardrail Name	Category	Conditions	Last Modified	Enabled
Critical	Warn on Category Data Leak	Proprietary data leakage		10 Dec 2025 10:48am UTC	☒
High	Grok - Sensitive Data Reduction	Privacy data sharing violations		26 Nov 2025 11:21pm UTC	☒
Critical	Block Finance Category	Privacy data sharing violations		20 Nov 2025 06:31pm UTC	☒
Critical	Block MS copilot	Proprietary data leakage		18 Nov 2025 09:57pm UTC	☒
Critical	Block PI, PHI and PCI data leakage	Proprietary data leakage		05 Nov 2025 04:03pm UTC	☒
Medium	Redact sensitive leak to Gemini	Privacy data sharing violations		04 Nov 2025 07:16pm UTC	☒
High	Gemini - Sensitive Data Reduction	Privacy data sharing violations		29 Oct 2025 10:15pm UTC	☒

Visibility You Can Act On

AI governance isn't a "set it and forget it" project. To find out if your guardrails are doing their job, you need clear answers to some very practical questions:

- What sensitive data did we successfully stop from being shared?
- Where did we accidentally block data that should've been allowed?
- What slipped through the cracks and made it into an AI prompt anyway?
- Are there insiders (well-meaning or otherwise) who keep testing the limits?



Semantic DLP makes those answers easy to find. High-level dashboards show you which AI applications and users pose the greatest risk, while intuitive visuals reveal exactly how sensitive data is flowing to AI tools across departments. You can spot patterns, pressure points, and problems at a glance.

Need to go deeper? No problem. Every policy violation is fully logged — who did it, which rule was broken, what data was involved, which AI tool it was shared with, and precisely when it happened. And for when things really need a microscope, Semantic DLP records every prompt and every response, helping you demonstrate regulation compliance, speed investigations, and fine-tune policies before small gaps turn into big problems.

In short: fewer blind spots, faster answers, and AI governance you can actually trust.

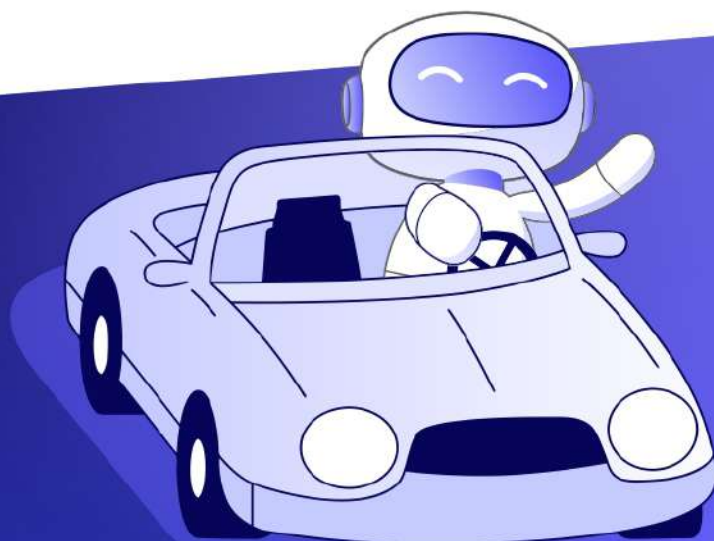


With Semantic DLP you can say yes to productivity without saying goodbye to your sensitive data or compliance:

- **Enable AI-driven innovation without increasing data risk:** Teams can use AI freely and productively, while sensitive data stays protected — no heavy-handed lockdowns required.
- **Eliminate blind spots caused by shadow AI:** Gain full visibility into every public AI tool in use, including unsanctioned and emerging applications, so nothing risky flies under the radar.
- **Apply precise, context-aware controls that users can't bypass:** Enforce AI policies based on what data means, not just keywords — resulting in fewer false positives and smarter enforcement.
- **Prevent sensitive data from being amplified by AI models:** Stop confidential information before it ever enters AI prompts, uploads, or responses — reducing downstream exposure and long-term risk.
- **Automatically protect new data the moment it's created:** Newly generated AI content is continuously discovered, classified, and secured without manual labeling or policy lag.
- **Gain actionable insight into AI risk and user behavior:** Identify your riskiest AI tools, users, and data flows with dashboards and detailed audit trails that support fast decisions and investigations.

Security That Doesn't Slow You Down

You don't have to choose between innovation and security. With the right data security governance in place, you can protect your data while empowering your teams to do their best work.



About Concentric AI

Concentric AI is intelligent data security made easy, helping businesses discover, understand, and protect sensitive information across cloud and on-prem environments. Its AI-powered data security governance solutions deliver precise visibility, automated protection, and safe AI adoption — so organizations can reduce risk, strengthen compliance, and confidently innovate without compromising what matters most.

Contact us today to see how we can help your organization enable AI without expanding your threat surface.

[Schedule a Demo](#)

