

BLACK HAT BUYER'S GUIDE

5 questions to ask every AI security vendor

A practical checklist for separating real AI protection from web-security features repackaged for the AI era.

1) Can you stop sensitive data before it reaches the AI service? “Do you inspect prompts, clipboard content, file uploads and desktop AI activity before the data leaves the endpoint?” Follow up: “Does your protection work only for browser traffic, or also for native apps, coding assistants and local agents?” What this exposes: Products that inspect only after traffic reaches a cloud proxy, when the sensitive data has already left the device. A strong answer should include: request-side inspection, endpoint enforcement, prompt and upload controls, and clear handling of encrypted or certificate-pinned applications.	2) How do you secure AI that never passes through your cloud? “How do you discover and govern local AI assistants, desktop applications, browser extensions and autonomous agents that do not traverse your service?” Follow up: “What protection remains when the user is offline or connected to an untrusted network?” What this exposes: Cloud-only architectures that equate AI security with categorizing websites or proxying web sessions. A strong answer should include: endpoint-native visibility and controls that continue to operate without depending entirely on cloud inspection.
3) Can you show me every AI tool and agent in use? “If we deployed your platform today, what sanctioned and unsanctioned AI activity would we discover tomorrow?” Follow up: “Can you distinguish a web app, desktop client, browser extension, API call and local agent — or do they all appear as a domain?” What this exposes: Visibility based mainly on URL categories, known application lists or traffic that already follows a managed route. A strong answer should include: actionable user, device, application and agent-level visibility, with policy controls that go beyond discovery dashboards.	4) How do you apply zero trust to AI agents accessing company data? “How do you authenticate, authorize and continuously control an AI agent accessing SaaS, private applications, APIs and enterprise data?” Follow up: “Can policies restrict the agent to the minimum data and services required, and can access be revoked immediately?” What this exposes: Platforms focused only on employee use of public chatbots, with limited control over agent-to-application and agent-to-data access. A strong answer should include: identity, device or workload posture, least-privilege access, encrypted connectivity, continuous policy enforcement and auditability.
5) What does your product replace, integrate with or conflict with? “Show me how your platform coexists with our current SWG, DLP, EDR, browser security and ZTNA agents.” Follow up: “Who performs TLS inspection, how is traffic routed, and what happens if two endpoint agents try to control the same session?” What this exposes: Hidden deployment complexity, double decryption, traffic-steering conflicts, excessive endpoint agents and unexpectedly high total cost of ownership. A strong answer should include: a clear deployment diagram, tested coexistence options, explicit exclusions, realistic licensing and a proof-of-value plan using your actual security stack.	

The question behind all five questions

Is this a purpose-built AI security platform, or an existing web-security product with AI labels added?
Ask vendors to demonstrate the difference using your users, your agents, your applications and your data.

Prepared by Veraify powered by Cloudbrink

Veraify secures and governs enterprise AI—including public tools, desktop apps, and autonomous agents—using endpoint-aware data protection and zero-trust access control. [Learn more at Veraify.ai/contact](https://veraify.ai/contact).