



Nightfall AI™

Protecting Sensitive Data from Shadow AI

**The Security
Practitioner's Guide**

Table of Contents

- I.** Introduction & Key Stats
- II.** What is Shadow AI?
- III.** Mapping the Risks of AI Apps
- IV.** Case Study: Disney NullBulge Hack
- V.** Identifying Shadow AI Activities
- VI.** Building a Governance Framework
- VII.** Securing Data with DLP
- VIII.** Step-by-Step Implementation Roadmap
- IX.** Conclusion & Next Steps
- X.** Appendix



I. Introduction

Artificial intelligence (AI) chatbots and other AI-driven tools have exploded in popularity, becoming go-to resources for employees looking to automate tasks, troubleshoot problems, generate content, and more. However, when these AI tools are deployed—or accessed—without direct oversight from IT or security teams, organizations risk introducing what is often called Shadow AI. Shadow AI refers to AI-powered applications, platforms, or plugins adopted by individuals or teams without formal security and compliance vetting. These tools can also include unsanctioned integrations with mainstream chatbots or generative language models.

Because AI chatbots are so easy to use—often requiring only a web browser or basic API key—employees may start interacting with them for everything from drafting emails to summarizing confidential files. They may unknowingly feed proprietary source code, sensitive customer data, or regulated personal information (PII) into these external models. Once that data leaves a controlled environment, it can be logged, stored, or even used to train external AI models. This leads to a significant risk of data leakage, breaches, or compliance violations under frameworks such as HIPAA, PCI DSS, and GDPR.



This guide presents a detailed playbook for security leaders who need to curb the risks of Shadow AI. We'll provide:

- A clear understanding of how Shadow AI emerges in modern workplaces.
- In-depth discussions of the primary categories of risk—ranging from compliance pitfalls to advanced cyberattacks.
- Best practices for identifying and governing chatbot use, including recommended policies and training programs.
- Techniques for applying data loss prevention (DLP), AI-based classification, and data lineage tracking to safeguard sensitive data.
- A phased implementation roadmap that organizations can follow to reduce their vulnerability to Shadow AI threats.
- A hypothetical scenario illustrating how a data leakage incident can occur—and be mitigated—within a single company's environment.

By adopting the strategies in this guide, CISOs and their teams can strike the right balance between empowering employees with AI-driven productivity tools and ensuring that regulated or high-value data doesn't leak out of the organization's control.



Key Stats

91%

Organizations using AI tools in some capacity (2025, Gartner).

68%

Employees that admit using unsanctioned AI tools at work (2024, McKinsey).

4.45M

Average cost per data breach globally (2024, IBM Cost of a Data Breach Report).

70%

Reduction in data leakage incidents with endpoint DLP solutions (2023, Forrester).

56%

Employees that are unaware of their organization's AI usage policies (2024, Ponemon Institute).

47%

AI vendors that retain user input for over 12 months (2024, TechCrunch).

54%

Employees that report a 20%+ productivity boost using generative AI tools (2024, McKinsey).

II. What is Shadow AI?

Defining Shadow AI

Shadow AI is any artificial intelligence solution—particularly chatbots or large language models (LLMs)—used within an organization without the knowledge or approval of the IT department, security teams, or compliance officers. It is the natural extension of “shadow IT,” where employees provision software or cloud resources informally. The difference here is that the data being fed to these AI models is often more sensitive, nuanced, or crucial than what might be uploaded to a typical unauthorized software platform.



Why It Happens

Shadow AI frequently arises because of:

- **Ease of Accessibility:** Many AI chatbot providers allow instant sign-up with a corporate or personal email, requiring no formal vendor onboarding or security review.
- **Employee Enthusiasm:** Tools like generative chatbots or code generation services can dramatically speed up tasks. Employees, especially those under deadline pressure, see immediate benefits and are often unaware of data protection implications.
- **Gaps in Official Guidance:** Where organizations do not proactively provide sanctioned AI tools or usage guidelines, employees fill that gap themselves—usually unaware of regulatory or contractual constraints.
- **Desire for Experimentation:** Tech-savvy staff might want to try emerging AI functionalities before official procurement or adoption processes catch up.

Common AI Tools Contributing to Shadow AI

- **Generative Chatbots:** Popular large language model-based chatbots that can write, summarize, or reason about user-submitted text.
- **AI Code Helpers:** Tools that assist in generating software code from prompts or scanning code for errors. Often these can store code snippets on external servers.
- **Image or Document Analyzers:** AI services that extract data from PDFs or images for quick indexing or transformation.
- **Browser Extensions:** Third-party plugins that embed AI functionality into everyday tasks (e.g., summarizing emails or chat threads).

In each instance, an employee may inadvertently share proprietary, regulated, or confidential information. If that data remains stored in the chatbot provider's infrastructure, it falls outside direct organizational control.



III. Mapping the Risks of AI Apps

This section examines the most critical risks introduced by employees feeding sensitive data into unauthorized AI chatbots or generative models.

Data Privacy & Regulatory Compliance

Potential Violations

- **HIPAA:** If an employee at a healthcare company uses an AI chatbot to summarize patient notes or medical histories, they could inadvertently share protected health information (PHI) with an external, non-HIPAA-compliant entity.
- **GDPR:** European Union data protection laws forbid transferring personal data to external processors without proper safeguards and consent, which might be lacking in a third-party AI provider.
- **PCI DSS:** Storing or transmitting credit card data in unauthorized systems is a direct breach of Payment Card Industry rules.

Violations can trigger hefty regulatory fines, along with reputational damage. Even if the organization can demonstrate no malicious intent, the compliance burden remains strict.

Difficulty of Enforcement

Data privacy regulations typically require robust data inventory, consent tracking, and strict transfer controls. Chatbots that rely on user inputs are unpredictable—employees might copy-paste personal details into an AI prompt. Without a formal control system, the organization cannot accurately track every piece of data leaving its environment.

Intellectual Property Concerns

Proprietary Data Leaks

Companies invest heavily in research, designs, source code, or special processes. If a user enters such proprietary content into an AI platform, that data could:

- Remain on external servers, accessible to the AI provider's training or internal teams.
- Resurface as "learned knowledge" in future chatbot outputs for other customers.

Competitive Espionage

Rivals could theoretically glean insights if the AI platform experiences a breach or decides to incorporate user-submitted data in a shared model. In industries like pharmaceuticals, financial services, or technology, losing such an edge can be devastating.

Model Retention Policies

Data Storage in Chatbot Systems

Many AI platforms log and retain user queries to refine their models. These queries can become part of the platform's training dataset. Even if data is stored in masked or tokenized form, advanced data correlation techniques or future system vulnerabilities might allow re-identification or partial reassembly of sensitive data.

Unclear Terms of Service

Often, the fine print states that data fed into an AI tool becomes part of its knowledge base or may be stored indefinitely. Even if employees "accept" these terms, it doesn't absolve the organization of responsibility for where the data ends up.

Cyberattack Vectors

Endpoints & Browser Extensions

Employees installing unauthorized AI chatbots or plugins may introduce new vulnerabilities. A malicious extension or Trojan AI tool could harvest credentials or act as a key logger, bridging the gap between internal systems and attackers.

Model Exploits & Prompt Injection

Advanced adversaries might manipulate AI inputs to reveal hidden data or system instructions. For instance, a compromised chatbot might inadvertently share an organization's previous inputs (including sensitive info) with a clever prompt that bypasses normal constraints.

Supply Chain Attacks

A third-party AI provider may have robust security, but what if a linked sub-service or plugin is compromised? Supply chain vulnerabilities often give attackers a path to data from unsuspecting users feeding proprietary information into the platform.



IV. Case Study: The Disney NullBulge Hack

The Disney NullBulge hack serves as a cautionary tale about the dangers of unauthorized AI tools, or "shadow AI," in corporate environments. This case highlights how employee use of unvetted AI applications can lead to severe security breaches, data exposure, and reputational damage.

In July 2024, Disney suffered a massive data breach when the hacktivist group NullBulge accessed and leaked 1.1 terabytes of sensitive data from the company's Slack channels. The breach exposed internal communications, source code, unreleased project details, financial data, and even personal employee information. NullBulge justified its actions by criticizing Disney's use of AI in content creation and its handling of artist contracts.

The hackers reportedly gained access through an insider vulnerability. Matthew Van Andel, a Disney engineer, unknowingly installed a malicious AI extension on his personal device. This malware, disguised as an AI tool or game mod, allowed NullBulge to infiltrate Disney's internal systems.



Key Risks Highlighted

- **Data Privacy Violations:** The breach involved sensitive employee data, potentially violating regulations like GDPR or CCPA.
- **Intellectual Property Exposure:** Proprietary source code and unreleased project details were leaked.
- **Cyberattack Vectors:** The incident underscores how shadow AI tools can act as Trojan horses. Malware embedded in an AI extension exploited Van Andel's credentials to access Disney's Slack.
- **Model Retention Policies:** Many generative AI tools retain user inputs for training purposes. If Disney employees had used unauthorized AI platforms for work-related tasks, sensitive data could be stored indefinitely on third-party servers.

Lessons Learned

Governance Over Shadow AI: Organizations must implement robust governance frameworks to monitor and regulate employee use of AI tools. This includes vetting all third-party applications for security risks before deployment.

Comprehensive DLP Deployment: Implement SaaS DLP solutions to scan and monitor collaboration platforms like Slack to reduce that attack surface if these apps are compromised. Complement SaaS DLP with endpoint and browser DLP, which can monitor where sensitive data is flowing (e.g., into unauthorized AI apps) and block such actions in real-time.

Employee Awareness & Training: Employees should be educated about the risks of shadow AI and the importance of using IT-approved tools only. Van Andel's case illustrates how lack of awareness can lead to catastrophic consequences.

Endpoint Security: Companies should enforce strict endpoint security policies to prevent unauthorized software installations. Advanced monitoring tools can detect unusual activity or unauthorized downloads in real-time.

AI-Specific Policies: Clear policies around AI usage are essential to mitigate risks associated with model retention and data sharing practices. Organizations should ensure employees understand the implications of feeding sensitive information into generative AI platforms.

V. Identifying Shadow AI Activities

Spotting and curbing Shadow AI usage is challenging; employees might not realize they're doing anything dangerous. Here are strategies for detection:

Symptoms & Red Flags

- 1. Unexplained Traffic:** Spikes in network traffic to known AI platform domains or suspicious new endpoints.
- 2. User Behavior Deviations:** Individuals who frequently copy large code snippets or text from internal repositories to external websites.
- 3. Unrecognized Subscriptions:** Credit card expense logs or SaaS expense statements listing AI services not in your official vendor list.

Technology Insights

Implementing monitoring and DLP solutions that recognize patterns of potential chatbot usage can illuminate hidden channels. For instance, a DLP engine might detect suspicious data exfiltration attempts if it sees multiple lines of personal data being transmitted to a known AI domain.

Employee Polling & Surveys

Since Shadow AI is often used by well-intentioned staff seeking productivity boosts, a non-punitive approach to discovering usage can help:

- **Anonymous Surveys:** Ask which AI tools employees currently use to do their jobs and in what capacity.
- **Focus Groups:** Encourage staff to share beneficial AI tools they've found. This approach fosters collaboration rather than punishment.
- **Open Discussion:** Promote an environment where employees feel comfortable disclosing their usage in the interest of security.



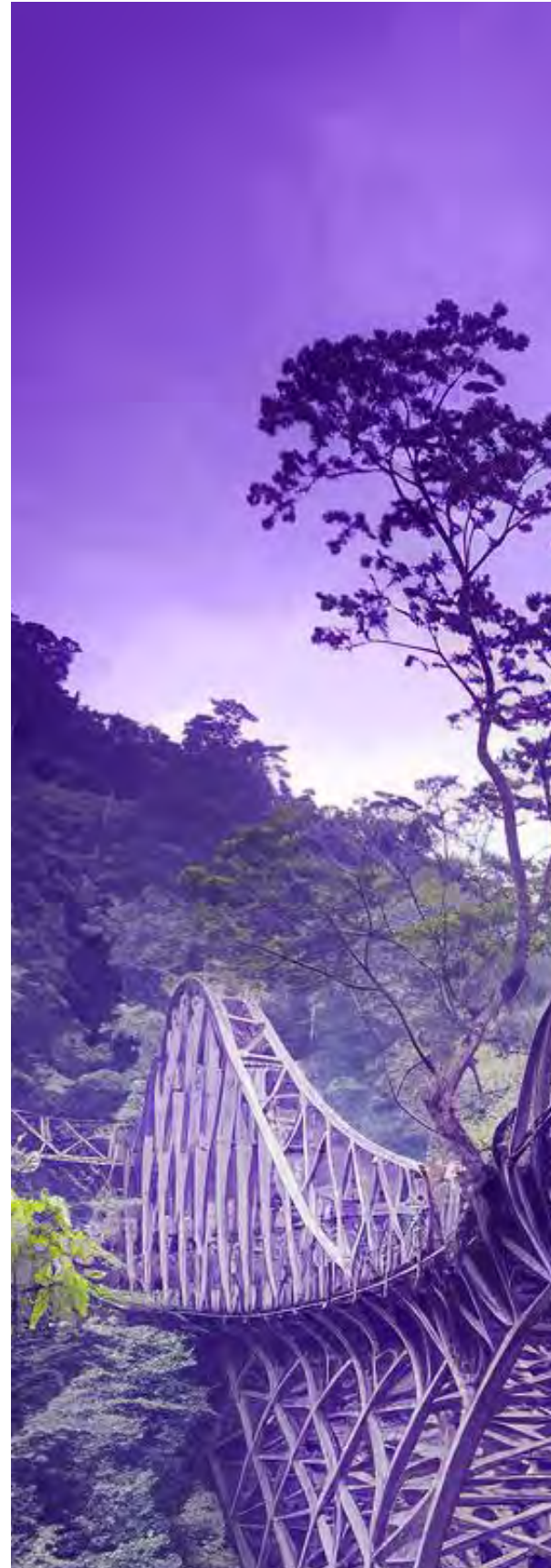
VI. Building a Governance Framework for Chatbot Use

Once you know your organization is exposed to Shadow AI, the next step is creating a comprehensive governance framework ensuring that legitimate AI usage meets certain standards, while unsanctioned usage is minimized or prevented.

Clear AI Policies

Organizations should articulate:

- 1. Acceptable AI Use Cases:** Define which tasks and data categories are permissible for AI chatbot usage. For example, allow employees to use an approved chatbot for grammar checks on marketing copy, but explicitly forbid uploading personal health information or other sensitive data.
- 2. Prohibited Data:** List data types that must never be entered into external AI (credit card numbers, patient health records, API keys, etc.).
- 3. Review & Approval Processes:** If teams want to adopt a new AI tool, they should follow a security review process, ensuring any personal data is masked or minimized before usage.
- 4. Incident Reporting:** Outline how employees can escalate issues if they inadvertently share restricted data or suspect a breach.



VI. Building a Governance Framework for Chatbot Use

Risk-Based Approvals

- 1. Tiered Risk Levels:** Tools that process no sensitive data might have fewer restrictions than those that analyze or store regulated PII.
- 2. Vendor Vetting:** Evaluate the data handling practices, compliance certificates, and data residency policies of AI tool providers.
- 3. Controlled Environments:** For internal use cases, an organization might spin up a private instance or on-prem large language model that does not feed data into a public cloud.



Employee Training & Awareness

- 1. Mandatory Workshops:** Conduct sessions educating employees on data classification levels, examples of real-life leaks, and how chatbots store data.
- 2. Phishing & Social Engineering:** Illustrate how malicious chatbots or plugins can trick staff into disclosing login credentials.
- 3. Safe Experimentation:** Provide a secure sandbox or clearly approved AI environment so users can experiment without risking production data.

VII. Securing Data with DLP

Implementing Data Loss Prevention (DLP) solutions that combine AI-based data classification and data lineage capabilities offers a powerful safeguard against Shadow AI. This section dives deeper into how these technologies work and the value they provide.

Endpoint & Browser-Based DLP for AI App Detection

Modern DLP solutions can deploy browser extensions and endpoint agents to monitor and control data as it flows between local devices and external services (including AI chatbots). By running directly on the user's machine and within their web browser, DLP gains real-time visibility over:

- Clipboard & Text Entries
- Network & Data Flows
- File Uploads & Attachments
- Audit & Investigation



Clipboard & Text Entries

Automated Classification: As soon as a user pastes or submits sensitive information—like PII or API keys—the endpoint agent or browser extension detects it.

Policy-Based Alerts & Blocking: Policies can either prompt the user with a coaching message (“You’re about to share sensitive data—are you sure?”) or block the action entirely if it violates security rules.



Network & Data Flows

Comprehensive Monitoring: Endpoint agents track traffic going to external domains. If large volumes of data, or specifically flagged data types, flow to suspicious or unapproved destinations, the system can intervene.

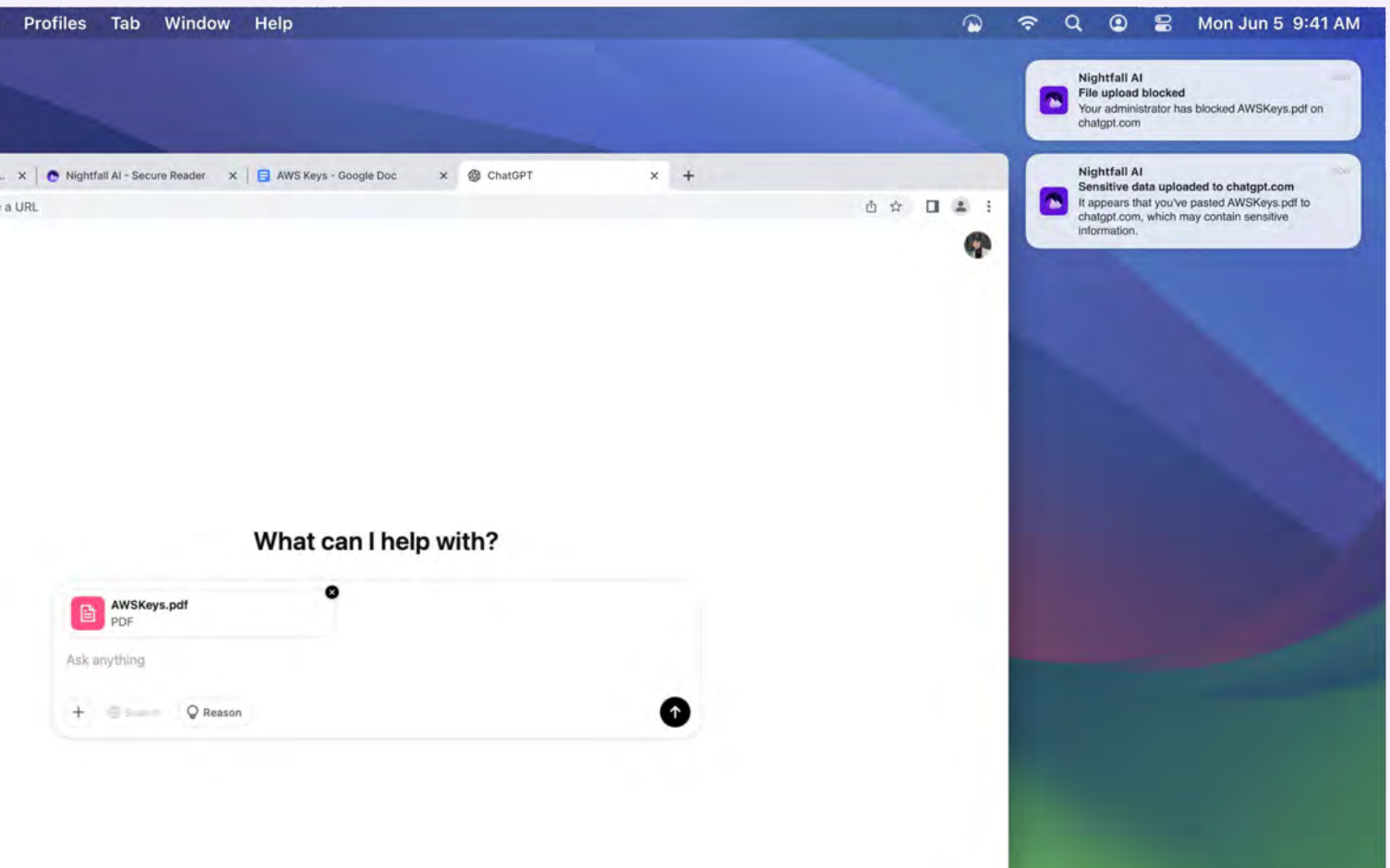
Thick Client Support: Beyond browser traffic, DLP at the endpoint can watch for data exfiltration attempts through desktop applications or command-line tools, stopping advanced leaks that wouldn’t be visible to network-level or API-based solutions alone.



File Uploads & Attachments

Scanning & Classification: Before the file leaves the device, the DLP agent checks attachments (code files, PDFs, images) for sensitive data.

Real-Time Interception: If a user attempts to upload regulated content (e.g., personal health information) to an AI chatbot website or other unauthorized cloud service, the DLP can alert the security team or stop the upload.



Audit & Investigation

Context-Rich Logging: Every attempted paste, upload, or transfer is logged with details about the data type, the user, the application used, and the time of the event. This granular insight helps investigators quickly reconstruct incidents.

Compliance Reviews: If a breach does occur, logs provide an authoritative record of how the data was exposed, ensuring you can meet regulatory reporting requirements.

AI-Driven Data Classification

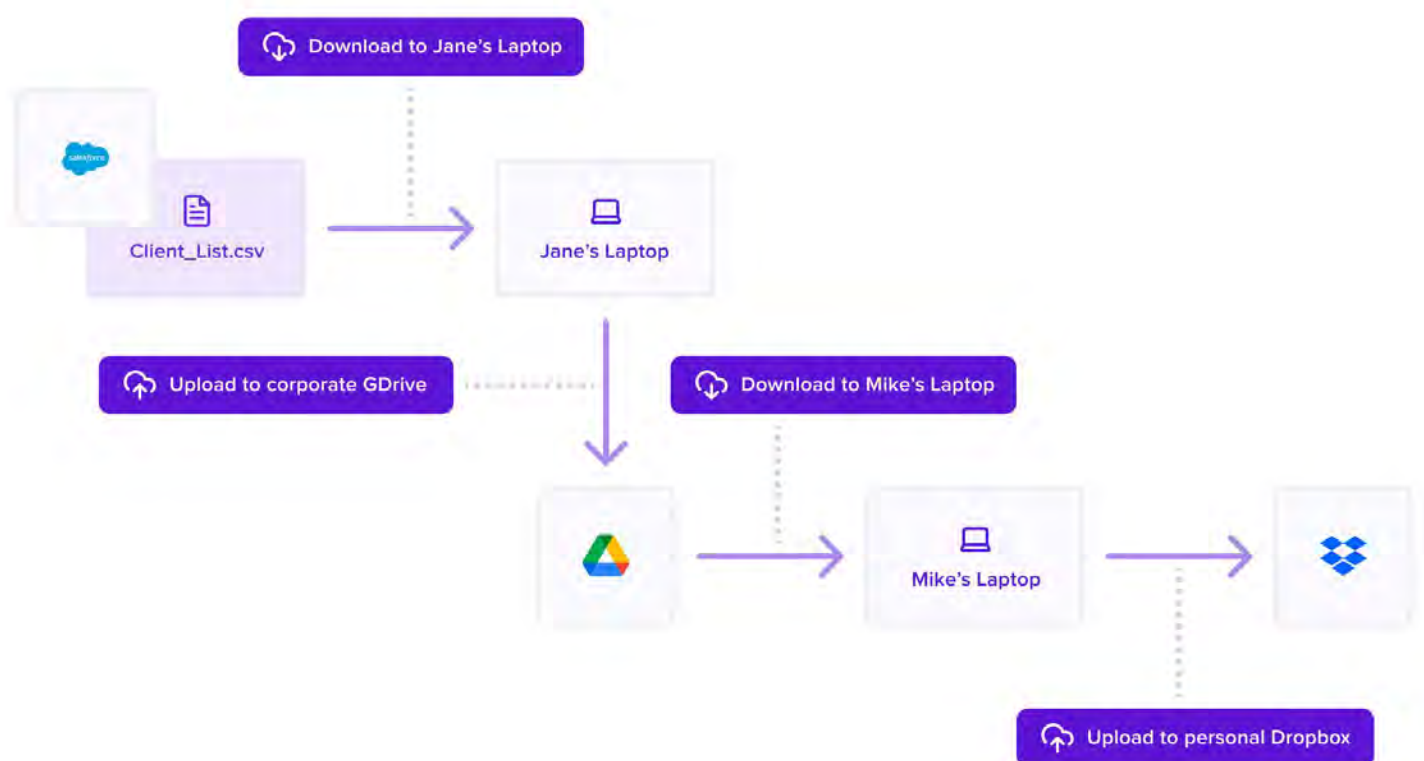
Traditional classification often relies on pattern matching (e.g., regex for credit card numbers). But modern AI classification can adapt to complex, context-heavy data:

- **Contextual Analysis:** AI models can identify sensitive data within paragraphs or code blocks, even if it's not a simple pattern (e.g., "John's social is 123-45-6789" vs. a disguised or partially masked variant).
- **Reduced False Positives:** Intelligent classification lowers the volume of meaningless alerts, ensuring security teams can focus on genuine threats.
- **Nuanced Language Detection:** If employees share sensitive details in foreign languages, specialized classification models can still detect potential violations.

Data Lineage Tracking

Data lineage addresses the question: "Where did this data come from, and where is it going?" By tracking lineage:

1. **Identification of Risky Flows:** If a file containing personal information migrates from a secure repository to a user's local device, and then to a known chatbot endpoint, lineage data shows the exact chain of custody.
2. **Faster Incident Response:** Investigations can pinpoint the user, device, and time of exfiltration, allowing for targeted remediation efforts.
3. **Proactive Policy Enforcement:** If a DLP system sees that data labeled "internal only" moves toward an external AI domain, it can instantly trigger blocking or encryption.



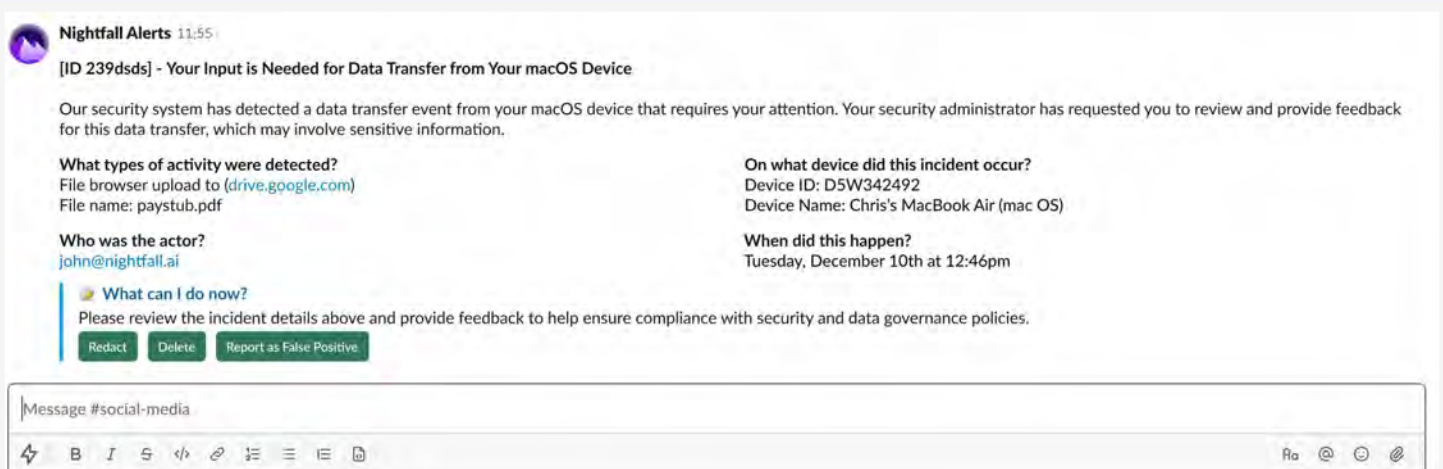
Automated Remediation

To handle real-time threats, an integrated system might:

- **Warn the User:** Notify the employee they are attempting to upload protected data into an unauthorized system.
- **Sanitize the Data:** Automatically redact or strip out sensitive fields.
- **Encrypt or Quarantine:** Temporarily lock down the file to prevent further exposure if malicious use is suspected.
- **Escalate to Security:** Provide immediate alerts to the SOC (Security Operations Center) for manual follow-up.

Additional Considerations

- **User Coaching:** A balanced approach often involves interactive prompts reminding employees not to share confidential or regulated information with unapproved AI tools. This is especially effective for preventing inadvertent disclosures while allowing legitimate work to proceed.
- **Integration with Other Systems:** Endpoint/browser-based DLP can feed logs into SIEM or SOAR platforms for broader incident correlation, or combine with network-level analysis for defense in depth.
- **Scalability & Performance:** Endpoint solutions must be lightweight and user-friendly to avoid interfering with productivity. Browser extensions, similarly, should be easy to deploy and automatically updated.



VIII. A Step-by-Step Implementation Roadmap

Establishing a robust defense against Shadow AI doesn't happen overnight. Below is a phased approach that many organizations find helpful:

Phase 1: Discovery and Inventory

- 1. Catalog Existing AI Tools:** Gather data on any known or suspected chatbot usage. Poll employees, check SaaS bills, and search for traffic patterns.
- 2. Identify Data Repositories:** Determine which systems contain sensitive or regulated data that could be at risk of inadvertent sharing.
- 3. Risk Assessment:** Conduct a preliminary risk assessment ranking each AI tool or usage pattern from low to high risk.



Phase 2: Policy Rollout and Employee Education

- 1. Draft AI Acceptable Use Policies:** Make them clear, jargon-free, and scenario-driven.
- 2. Obtain Executive Buy-In:** Ensure leadership endorses the policies so employees understand they're official and enforceable.
- 3. Awareness Campaign:** Host training sessions or lunch-and-learns. Provide real-world examples of data leaks to highlight the importance of compliance.

Phase 3: Deploy DLP & Monitoring Tools

- 1. Select a DLP Solution:** Opt for one with browser and endpoint coverage, as well as a robust set of API integrations into SaaS application for holistic coverage.
- 2. Configure Classification Rules:** Tailor detection to your sensitive data categories. Some organizations require specialized information types (e.g., for healthcare or financial data).
- 3. Run Test Audits:** Initiate pilot programs in a controlled environment. Encourage a small user group to attempt typical tasks with chatbots under watchful logging to calibrate the system.

Phase 4: Continuous Tuning and Incident Response

- 1. Refine Policies & Alerts:** After the pilot, reduce false positives and tighten critical rules.
- 2. Institute Regular Audits:** Evaluate logs monthly or quarterly to ensure compliance.
- 3. Respond to Incidents:** If an attempted exfiltration or inadvertent data share occurs, follow your defined incident response plan. Document lessons learned for ongoing improvement.

Phase 5: Ongoing Evaluation and Expansion

- 1. Add More AI Integrations:** As other business units adopt new AI tools (for marketing, HR, finance, etc.), bring them into the same governance framework.
- 2. Evolve Your Policies:** Data sensitivity definitions can change with new regulations or acquisitions. Keep guidelines updated.
- 3. Stay Alert to New Attack Vectors:** As AI technology advances, new exploit methods will arise—regularly consult threat intelligence feeds and adapt your defenses.



IX. Next Steps

Key Takeaways

- 1. Shadow AI is Real, Growing, and Risky:** The faster employees adopt AI chatbots without oversight, the greater the chance of accidental data spills or compliance violations.
- 2. No Single Tool or Policy Is Enough:** Successful risk mitigation depends on a combination of governance frameworks, user training, robust DLP technology, and continuous monitoring.
- 3. Data Lineage & AI Classification Are Game-Changers:** Being able to see both the “what” (content classification) and the “how” (lineage and movement) of data transforms detection and response capabilities.
- 4. Early Employee Engagement Is Critical:** Shadow AI typically begins with legitimate business goals—fostering an open dialog and providing safe, approved AI alternatives can preempt unauthorized usage.

Action Items

- **Publish or Update AI Policies:** Ensure they are easily accessible to every department.
- **Deploy or Upgrade DLP:** Look for advanced solutions with real-time scanning, AI-based classification, and data lineage mapping.
- **Launch Training Programs:** Host mandatory workshops that show real-world data leakage examples and clearly define the do’s and don’ts of AI chatbot usage.
- **Conduct Ongoing Audits:** Verify that new tools and workflows comply with your organization’s guidelines. Adjust policies and detection rules as the AI landscape evolves.

Strategic Importance

In a world where AI advancements outpace regulatory guidance, organizations that tackle Shadow AI head-on will build a reputation for robust data protection—instilling confidence among customers, regulators, and employees alike. Moreover, a strong approach to AI adoption can enhance productivity while minimizing legal, financial, and reputational risks.





Nightfall AITM

www.nightfall.ai

X. Appendix

Policy Templates

A sample “AI Acceptable Use Policy” might include:

- **Purpose:** “To ensure employees use AI services responsibly, protect regulated data, and maintain compliance.”
- **Scope:** “All employees, contractors, and third-party vendors must comply when handling sensitive data.”
- **Definitions:** Clarify “AI chatbot,” “sensitive data,” “restricted data,” etc.
- **Approved Tools:** E.g., “Only [organization-sanctioned AI platform] may be used for production data.”
- **Prohibited Actions:** E.g., “Do not upload HIPAA or PCI-regulated data to any external chatbot.”
- **Violations & Enforcement:** Consequences for repeated misuse or egregious negligence.

Assessment Framework for Shadow AI

Use a basic scoring system (e.g., Low, Medium, High) for each of the following:

- **Data Sensitivity:** Does the team handle regulated or proprietary data?
- **AI Tool Maturity:** Has the vendor undergone a formal security review?
- **Employee Training:** Are workers aware of acceptable use policies?
- **DLP Integration:** Is the AI tool integrated or recognized by your existing security solutions?

A high cumulative score might indicate a heightened need for stricter approval processes.

Glossary

- **AI Chatbot:** A system that processes natural language inputs and provides responses, often powered by large language models.
- **Generative AI:** AI capable of producing new content (text, images, code) rather than merely analyzing or summarizing existing data.
- **Shadow IT:** Any software or cloud resource used by employees without IT approval.
- **Prompt Injection:** A security vulnerability in which malicious instructions are hidden in user input to manipulate an AI system’s behavior.
- **LLM (Large Language Model):** A category of AI models trained on massive text corpuses that can generate human-like text.