

AI and the Expanding Attack Surface: How Generative Tool Adoption Creates New Risk Vectors in Mid-Market Businesses

Bridgham Technology Institute

Bridgham Cyber, LLC

June 2026

Reissued September 2026 under the Bridgham name.

Abstract

Generative AI adoption among mid-market veterinary practices and insurance agencies accelerated materially in 2026, driven primarily by productivity applications in clinical documentation, policy quoting, client communication, and administrative workflow. This adoption occurred largely without corresponding security assessment, creating a new category of risk exposure that current compliance frameworks do not address and that most IT providers are not configured to monitor. The risk is not theoretical. It is operational: AI tools connected to core systems are ingesting sensitive data, generating outputs that may contain protected information, and operating on permission scopes that frequently exceed what the underlying workflow requires. The purpose of this analysis is not to discourage AI adoption, which delivers genuine operational value. It is to provide the framework for adopting these tools with a clear-eyed view of the access, data, and supply chain exposures they introduce.

1. Threat Environment

1.1. The Permission Scope Problem: AI Tools Requesting Broad System Access

The most immediate security concern associated with AI tool adoption in the mid-market is not adversarial AI. It is the routine permission scope requested by productivity AI applications when integrated with core operational systems. AI writing assistants integrated with email clients frequently request access to the entire mailbox, including sent items, attachments, and historically received messages. AI scheduling tools integrated with practice management systems may request read access to patient records in order to contextualize appointment scheduling. The individual user authorizing these integrations typically does so without reviewing the permission scope, because the workflow benefit is immediate and the security implication is not.

The aggregate effect of multiple AI tool integrations across a 30-person practice or agency is a significant expansion of the data access surface. Each tool represents a potential point of data exfiltration, either through a security incident at the AI vendor or through the AI system's data retention and training practices, which vary considerably across providers and are frequently updated without prominent notification.

1.2. Shadow AI: Staff-Initiated Tool Adoption Outside IT Governance

Shadow IT, the use of software tools not approved or inventoried by organizational IT governance, is not a new phenomenon. Shadow AI is its current manifestation, and it carries a materially higher data sensitivity profile than prior shadow IT categories. When a staff member uses a personal ChatGPT account to draft a client letter, pastes client financial information into a generative AI interface to summarize a policy document, or uses a consumer AI tool to process medical record data for a clinical note, the data transferred to that tool is subject to the terms of service and data handling practices of the AI provider, not the organization's security program.

Gartner research published in 2025 indicates that more than 55 percent of employees at organizations without a formal AI use policy report having used a personal AI tool for work-related tasks. For the average mid-market insurance agency or veterinary practice, this means that a majority of staff are likely using AI tools for work tasks that the organization has not reviewed, approved, or inventoried.

1.3. AI-Enhanced Social Engineering: Elevated Phishing Quality at Scale

Generative AI has materially reduced the cost and increased the quality of phishing and social engineering attacks targeting mid-market businesses. Prior to widespread AI availability, high-quality targeted phishing was resource-intensive, limiting its application to high-value targets. Generative AI enables personalized, grammatically precise, contextually accurate phishing content at near-zero marginal cost.

The FBI IC3 2026 report notes a significant increase in reported business email compromise incidents where the initial phishing content was assessed to have been AI-generated based on linguistic analysis. For insurance agencies processing high-value financial transactions and veterinary practices receiving payments for specialty procedures, elevated phishing quality represents a direct increase in BEC and wire fraud exposure that existing security awareness training, designed around prior-generation phishing characteristics, does not fully address.

Key Observation. *AI adoption in mid-market businesses is primarily a productivity story, and appropriately so. The security dimension is not a reason to avoid adoption. It is a reason to govern adoption with the same rigor applied to any other vendor relationship that touches sensitive data.*

2. Regulatory Update

2.1. FTC AI Guidance and Section 5 Applicability

The FTC has issued guidance indicating that AI tools used in consumer-facing contexts, including insurance recommendations and health-related communications, are subject to Section 5 of the FTC Act's prohibition on unfair or deceptive acts. While the guidance does not create AI-specific compliance obligations for most mid-market businesses, it signals that the use of AI tools in client-facing workflows without adequate oversight creates regulatory exposure if the AI output is

inaccurate, biased, or results in consumer harm.

For insurance agencies using AI tools for preliminary policy recommendations or underwriting summaries, and for veterinary practices using AI for client communication about clinical findings, the practical implication is that organizational review processes for AI-generated client-facing output should be documented. The FTC’s posture suggests that the question in a future enforcement action will be: what did the organization do to review AI output before presenting it to consumers.

2.2. HIPAA and AI-Generated Clinical Documentation

The HHS Office for Civil Rights has not issued AI-specific HIPAA guidance as of publication date. However, OCR’s existing guidance on cloud computing and third-party service providers applies to AI tools that process protected health information. A veterinary practice using an AI transcription or clinical documentation tool that processes patient data is operating in a framework where BAA obligations apply if the data meets the PHI threshold under HIPAA.

The practical implication for veterinary practices exploring AI clinical documentation tools is that vendor selection should include a review of the vendor’s willingness to execute a BAA and their ability to demonstrate technical safeguards commensurate with the sensitivity of data processed. AI vendors that are unwilling to execute a BAA for tools processing PHI-adjacent data should not be deployed in clinical workflows without legal review of the applicable data category.

2.3. North Carolina Consumer Protection and AI-Driven Communications

The North Carolina Consumer Protection Act’s prohibition on unfair and deceptive trade practices applies to business communications generated or assisted by AI tools in the same manner it applies to other business communications. There is currently no NC-specific AI regulation in effect, but the AG’s office has referenced federal AI guidance in consumer complaint resolutions, suggesting that NC enforcement posture will track federal agency guidance as it develops.

3. Intelligence Briefing

3.1. Deep Dive: An AI Governance Framework for the Mid-Market Practice and Agency

AI governance in the mid-market context does not require a dedicated AI ethics committee or a formal model risk management function. It requires three operational capabilities: knowing what AI tools are in use and what data they access, establishing a written policy that governs permissible use, and applying the same vendor oversight framework described in Issue 4 to AI tool vendors.

Capability	Description and Implementation
Capability 1: AI Tool Inventory	A complete list of AI tools in use across the organization, including consumer tools used by staff without IT approval. The inventory should include the data categories each tool accesses, the permission scope granted, and whether a data processing agreement or BAA is in place where required. This inventory follows the same methodology as the vendor inventory described in Vigil Issue 4.

Capability	Description and Implementation
Capability 2: AI Use Policy	A written policy governing permissible AI tool use by staff. The policy should specify: approved tools and workflows, prohibited uses (including the input of client, patient, or financial data into non-approved consumer AI tools), the review requirement for AI-generated client-facing output, and the process for requesting approval of new AI tools. The policy requires no technical implementation. It requires communication and acknowledgment.
Capability 3: Vendor Security Review	AI vendors holding access to sensitive data should be subject to the same Tier 1 vendor oversight described in Vigil Issue 4: security documentation review, data processing agreement or BAA, and defined incident notification obligation. The specific additional questions relevant to AI vendors include: data retention practices, whether the vendor uses customer data for model training, and how the vendor handles data access requests or deletion requests.

The AI governance framework described above requires approximately 12 to 16 hours of initial implementation effort for an organization at the mid-market scale. The inventory typically surfaces three to five AI tool relationships that the organization had not previously categorized as security-relevant. The use policy provides the behavioral governance that prevents shadow AI accumulation from reconstituting the same exposure surface after the initial inventory has been completed.

The most operationally significant finding from AI governance engagements conducted in our client portfolio is not the AI tool that has been deployed with appropriate oversight. It is consistently the consumer AI tool, often a free-tier subscription held by an individual staff member, that has been processing the organization's most sensitive data for months or years without any institutional awareness that it existed.

4. Recommendations

- 1. Conducting an AI tool inventory and applying Tier 1 vendor oversight to AI vendors holding sensitive data.** Clients are treating AI tool governance as an extension of the vendor tiering model introduced in Issue 4, not as a separate initiative. Generative AI tools with access to client, patient, or financial data are classified as Tier 1 vendors regardless of price point or perceived category. This classification triggers the security documentation review and contractual notification requirements applied to all Tier 1 vendors. The operational lift is minimal because the methodology is already in place.
- 2. Issuing a written AI use policy and obtaining staff acknowledgment.** Clients have drafted and distributed a one to two page AI use policy that defines approved tools, prohibited data inputs, and the review requirement for AI-generated client-facing output. The policy is communicated in a staff meeting rather than distributed as a document only, ensuring that the behavioral expectations are understood and not merely documented. Acknowledgment is recorded for compliance purposes.
- 3. Updating security awareness training to address AI-enhanced phishing.** Standard security awareness training scenarios use phishing examples that predate generative AI. Clients are updating training materials to reflect the current phishing environment, specifically the absence of the

grammatical and formatting cues that staff have been trained to identify as phishing indicators. The updated training emphasizes sender verification, link behavior, and out-of-band confirmation for high-risk requests, rather than content quality as the primary detection signal.

The Bridgham Technology Institute is the advanced research division of Bridgham Cyber for cybersecurity and artificial intelligence research. Issues 1 through 5 constitute the inaugural series. Future series will be announced to clients and distribution list subscribers through the Bridgham Cyber client portal.

Disclaimer

This publication is produced by Bridgham Cyber, LLC for informational and analytical purposes only. It reflects independent research and does not constitute legal, regulatory, compliance, or investment advice, nor a solicitation of any kind. The analysis draws on sources believed to be reliable at the time of writing, and no warranty is made as to their accuracy or completeness. Reproduction or redistribution, in whole or in part, requires attribution to Bridgham Cyber, LLC.