



AI Governance Is Not a Checkbox: Practical Lessons from an Interdisciplinary Panel

On September 17, I moderated a D.C. Bar Intellectual Property Law Community Artificial Intelligence Subcommittee panel, [Conducting an Artificial Intelligence/Cyber Security Audit – Assessments and Key Considerations](#). Artificial intelligence (AI) is embedded across nearly all industries; it is in our e-mails, pervasive in research, aiding in patient diagnosis and supporting clinical decisions. Yet, AI adoption often outpaces governance. As a dual-licensed registered nurse and attorney practicing at the intersection of health care, cybersecurity, privacy, and AI, I moderated a panel discussion that addressed both frontline operational and legal considerations.

The panel included complementary perspectives. Stephanie Spaulding, founder and principal advisor of [Synova](#), advises health care and life sciences organizations on AI governance, data infrastructure, and audit-ready system design, drawing on clinical laboratory experience and more than a decade of health informatics leadership. [Dr. Bishnu Sarker](#), assistant professor of health informatics, University of North Texas College of Artificial Intelligence and Advanced Analytics, develops AI and machine-learning models, including large language models, for life sciences, medicine, and health. Spencer Fane attorney [Shawn Tuma](#) has practiced in cybersecurity, data protection, and technology law since 1999 and represents businesses ranging from small and midsize companies to Fortune 100 organizations in AI, cybersecurity, privacy, incident response, compliance, computer fraud, and related litigation. Together, the panelists' legal, academic, and operational perspectives converged on one point: responsible AI use is not a one-time approval decision, but an ongoing risk-management discipline.

An AI Audit Is a Continuing Process

Dr. Sarker opened the panel with a presentation on why AI language models hallucinate, and how hallucination is baked into the objective. A language model's job is to assign probability to word sequences – $P(w_1 \dots w_n)$, or the next word given prior context. Nowhere in that objective is there a term for truth. The model optimizes for fluency and plausibility, so a confident, well-formed falsehood is exactly what a well-trained model is built to produce. Hallucinations are not a “bug” introduced by large models; they are the residue of an objective that rewards likely text. Mitigation has to come from outside the language-modeling objective, including retrieval from trusted sources, citation verification, and human review.

Providing an overview of the AI tool audit process, Spaulding emphasized that activating an AI tool begins a continuing relationship with the vendor handling the organization's data. A pre-deployment review is only a starting point. Vendors may change models, subprocessors, retention practices, controls, or features, while internal uses may expand beyond the original plan. Meaningful audits must revisit these variables periodically and after material changes.

Spaulding added that organizations should maintain an inventory of tools and uses, assign owners, document data flows, review vendor commitments, test controls, and preserve decision records. “Passing” an audit does not mean risk has disappeared; it means risks are known, assigned, monitored, and addressed within defined tolerances.

Shadow AI Is Often a Signal, Not Simply Misconduct

Spaulding then discussed how to approach employee use of “shadow AI.” Shadow AI refers to use of AI tools or features by employees without the approval, knowledge, or oversight of an organization's information technology group or security.

Organizations may assume that employees using unapproved AI tools are disregarding policy, Spaulding noted. She then offered a more useful diagnosis: most shadow AI reflects people looking for a faster way to complete genuine work. A purely punitive response, she explained, may drive usage further underground, leaving leadership with less visibility and employees with fewer safe options.

The better first step, Spaulding offered, is discovery: ask which tools employees use, for what tasks, and what needs approved systems fail to meet. Monitoring, expense data, surveys, and candid conversations can build a realistic inventory, she said. Policies should distinguish low-risk experimentation from uses involving prohibited confidential, privileged, personal, proprietary, or regulated information, Spaulding said.

Fluent Output Is Not the Same as Reliable Output

Dr. Sarker then explained why generative AI can appear remarkably capable while remaining inherently uncertain. Modern systems break language into tokens, translate these tokens into numerical representations, and calculate probable continuations. They are optimized to generate plausible responses but not to recognize when they lack sufficient information or to guarantee truth.

This technical reality produces two essential rules: protect the data entering the system and verify consequential outputs. Users should not provide personal, confidential, or sensitive information unless they understand how the system stores, uses, and shares it. Important factual, legal, clinical, financial, or operational outputs should be checked against reliable sources or reviewed by a qualified professional. AI is a powerful assistant, but it is not a substitute for human judgment. Once the information is fed into the model, it often cannot be pulled back out, Dr. Sarker noted.

Human review, however, is not a magic safeguard. Spaulding cautioned that reviewers who are rushed, overloaded, or overly trusting may merely rubber-stamp machine output. Organizations should define what reviewers must examine, give them enough time and subject-matter context, and escalate higher-risk uses to appropriately trained personnel, Spaulding explained. Review quality – not simply the presence of a person in the workflow – is what matters, she emphasized.

Cybersecurity and Data Governance Are Inseparable

Tuma underscored that cybersecurity and data are intertwined for every business. His practical yet memorable warning, “Data is the hot potato,” captured the need to understand the economic and contractual structure behind “free” tools. If an organization is not paying for a licensed version, it may effectively be the product. An

organization should not assume that a familiar interface, reputable brand, or slick vendor resolves questions about prompt retention, model training, account controls, incident notification, or downstream access, Tuma explained.

Vendor diligence should address what data the tool collects and where it is stored; whether prompts are used for training; who can access data; retention and deletion practices; subprocessors; logging and administrative controls; and responsibility for security incidents, intellectual property concerns, and inaccurate outputs. Technical and contractual protections should reinforce one another, Tuma added.

Production Use Changes the Risk

The panel emphasized that a polished demonstration is not proof that a system will perform reliably in daily operations. Real-world data may be incomplete, inconsistent, outdated, or materially different from test data, while workflows introduce edge cases and user behaviors that a pilot may reveal before deployment. Before scaling, organizations should test representative conditions, define acceptable error rates, confirm security and privacy controls, establish fallback procedures, and identify when use must stop.

Five Practical Steps to Take Now

In summary, the panel articulated the following steps for conducting an AI audit:

1. **Inventory actual use.** Identify tools, users, purposes, data types, integrations, and decision points – not merely approved applications.
2. **Classify risk by use case.** Apply stronger controls where AI affects rights, legal obligations, sensitive data, safety, finances, or external communications.
3. **Evaluate vendors continuously.** Review contracts, security documentation, retention practices, model changes, subprocessors, and administrative controls at onboarding and periodically thereafter.
4. **Design meaningful human review.** Specify reviewer qualifications, verification steps, escalation criteria, and documentation requirements.
5. **Create a reporting culture.** Give employees a clear, nonpunitive path to disclose AI use, errors, and suspected data exposure promptly.

The Bottom Line

The strongest lesson from the panel was that responsible AI governance must be practical, interdisciplinary, and continuous. Organizations do not need to eliminate experimentation, but they do need visibility into it. They do not need perfect certainty, but they do need controls proportionate to the stakes. And they should not treat human review, vendor assurances, or a completed audit as substitutes for sustained oversight.

AI adoption is ultimately a governance challenge as much as a technology decision. The organizations best positioned to benefit will be those that understand how their people actually use AI, protect the data that makes it valuable, verify the outputs that matter, and revisit their safeguards as tools and risks evolve.

This blog post was drafted by [Christine Chasse](#), an attorney in the Plano and Dallas, Texas, offices of Spencer Fane. For more information, visit www.spencerfane.com.

Click [here](#) to subscribe to Spencer Fane communications to ensure you receive timely updates like this directly in your inbox.