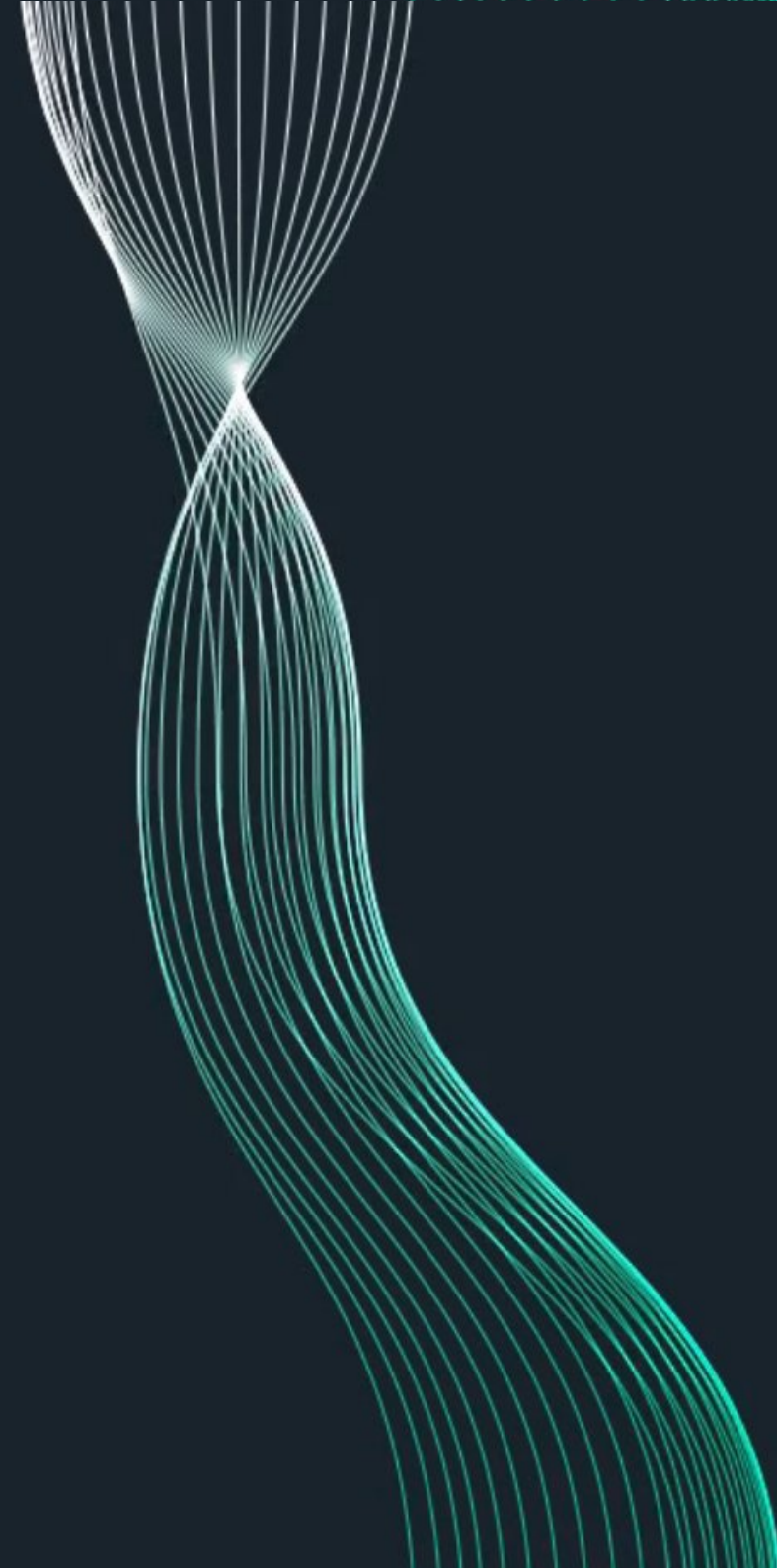


AI in Ethics & Compliance:

Risk to Manage, Tool to Leverage

ETHISPHERE[®]
GOOD. SMART. BUSINESS. PROFIT.[™]

 **SpeakUp**[®]



Report at a Glance

ETHISPHERE
GOOD. SMART. BUSINESS. PROFIT.



 Davis Wright
Tremain LLP

FOREWORD

AI GOVERNANCE + OVERSIGHT

AI IN PRACTICE

AI OUTCOMES FOR E&C

USING AI INSIDE THE E&C FUNCTION

EPILOGUE

Foreword

New Challenges, Current Skills

As we look forward to 2026, Ethisphere remains steadfast in our mission to advance business integrity. Artificial intelligence (AI) is no longer a future reality. It's a business imperative for ethics and compliance teams globally as they work to advance business integrity within their organization, and by extension, across the business itself.

At many companies, Ethics & Compliance (E&C) teams became early stakeholders in organization-wide AI as they created the governance around the use of this technology. As they did that, E&C got an unobstructed view of how transformative—and disruptive—AI can be. As the regulatory environment continues to evolve through legislation and litigation, teams will need to stay engaged with the changing expectations, just as they have for other risk areas.

Fortunately, no one is better positioned than E&C to do exactly that.

With so many E&C teams asked to deliver an impact beyond what their departmental budgets would suggest, the case for E&C's own use of AI has never been clearer. You are ideally situated to use this technology in a way that radically advances how much your team can actually do on behalf of your organization. This report will offer clear instruction and examples of how, as well as ideas of where we see the risks and opportunities evolving.

A stylized, handwritten-style logo for Erica Salmon Byrne in a light teal color.

**ERICA
SALMON BYRNE**
Chief Strategy Officer, Ethisphere



Executive Summary

E&C leaders are done debating whether or how to use AI. We're executing. 77% of E&C teams now hold a significant or coordinating role in AI governance, and 57% have employees trained in AI. Crucially, governance is the unlock for outcomes. When guardrails are codified into approvals, contracts, and monitoring, compliance can safely reap the rewards and implement AI with defensible oversight.

These things alone don't sufficiently manage AI risk, however. Most AI enters the organization via vendors, and while 84% of E&C teams own third-party risk, only 15% include AI provisions in third-party codes and just 14% have audited even half of their vendors.

This report closes that gap with insights into the AI regulatory landscape and governance framework examples, plus RACIs, model inventories, impact reviews, and HITL triggers. And since E&C leaders want less theory and more proof, we are also sharing peer stories from Ethisphere's BELA companies Verisk, Cargill, and Palo Alto Networks to showcase their literacy-first adoption and RAG-enabled workflows that shift time from drafting to judgment.

THE MANDATE IS CLEAR:
move from policy to proof,
extend governance downstream
to third parties, and convert
compliance's seat at the table
into measurable program
performance.

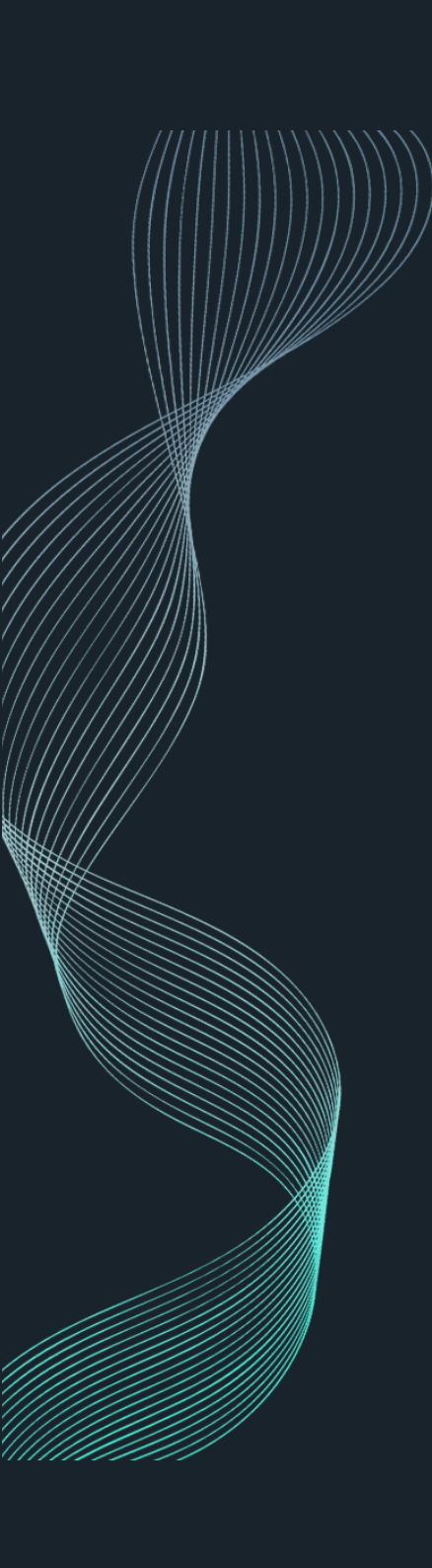
Tim
Morss

Chief Executive Officer, SpeakUp



AI Governance + Oversight

Expert Insights from Davis Wright Tremain



Artificial intelligence is no longer a niche technology. If it hasn't already, it will soon permeate nearly everything that our organizations, and we as Ethics & Compliance professionals, do. It impacts hiring platforms, enables personalized customer experiences, underpins our supply chains, and drives efficiency and innovation across organizations. But alongside this transformative potential comes substantial (and often novel) legal, ethical, and reputational risks. The challenges of managing these risks are increased by a patchwork regulatory landscape, rapidly changing technology, and the potential for significant harm if AI systems are misused.

The potential consequences of failing to properly govern AI and its use span far beyond regulatory risk. Imagine that your company loses control of the data that it planned to monetize as a central aspect of its business strategy. Imagine that your business spends extensive time and resources to develop an AI system that it then cannot commercialize or sell to key target clients. Imagine that your company is hit with a massive copyright lawsuit that brings not only monetary damages, but a threat to the viability of your AI system. These are not theoretical scenarios, and many ethics and compliance teams are, at best, playing catch up.

Whether your organization is developing or deploying AI systems, the mandate for ethics and compliance professionals is clear. We must apply our skills, cross-functional perspectives, and cross-organizational visibility to tackle the challenges that AI presents. Our ability to identify not only existing legal risks but also those on the horizon; to translate those risks into concrete, practical controls; and to design tailored training and sophisticated governance and monitoring frameworks are crucial to the development of effective AI governance and compliance.



TIP: Listen along to Bill Coffin as he narrates this section. Look for the audio cue in the upper right.

▶ [Click link to listen to audio](#)



**Davis Wright
Tremain LLP**

High-Level Overview of the Regulatory Landscape

The AI regulatory landscape is rapidly developing in numerous jurisdictions in the U.S. and across the globe. As discussed in more detail below, this landscape is a patchwork of state laws regulating various aspects of AI, a handful of U.S. federal rules and enforcement actions, and international legal regimes directed at AI. There also are numerous other broadly applicable legal regimes that apply to AI's many uses. Keep in mind that if it is not legal for you to do something, it is also not legal for AI to do that same thing on your behalf. Below, we review some of the key relevant compliance considerations, including privacy, data security, antitrust, copyright, and confidentiality, as well as provide additional considerations relevant to government contractors.

AI Governance and
Responsible AI

Antitrust

Data Privacy

Confidentiality

Data Security

Government Contractor-
Specific Considerations

Copyright

AI Governance & Responsible AI

ACROSS THE UNITED STATES

California laws require certain developers of generative AI systems to disclose what data was used to train such systems, to create AI detection tools that allow users to determine whether content is AI-generated, and to include “latent” disclosures in such content. California and other states have also enacted laws governing the performance of personal or professional services by AI-generated “digital replicas” of individuals. Utah has enacted laws governing AI chatbots that interact with consumers, imposing rigorous requirements on tools that offer mental health services and prohibiting the unauthorized commercial use of simulated or artificially recreated personal identities, including simulations created using generative AI. At least 38 states have enacted laws governing some aspect(s) of AI. Colorado is the first state thus far to enact a comprehensive law governing AI systems, and others are likely to follow. Among other things, the Colorado AI Act imposes substantial restrictions

and compliance obligations on developers and deployers of high-risk AI systems that are intended to interact with consumers and make or be a substantial factor in making “consequential decisions” in such areas as employment, insurance, housing, credit, education, and healthcare. Texas initially introduced legislation modeled on the Colorado AI Act but ultimately enacted a law that restricts the development and deployment of AI for a much narrower set of purposes, including the creation of child pornography and behavioral manipulation.



states have enacted laws governing some aspect(s) of AI.

THE BIG PICTURE

Although the U.S. Congress has not enacted a federal law to regulate the development or deployment of AI, the Trump administration has adopted an [AI Action Plan](#) designed to maintain U.S. leadership in AI by:

- Promoting AI exports
- Streamlining the development of infrastructure that supports AI
- Mandating “ideological neutrality” in AI models

In addition, federal agencies like the National Institute of Standards and Technology (NIST) have developed [AI governance frameworks](#) for companies to use to mitigate and manage AI risks.

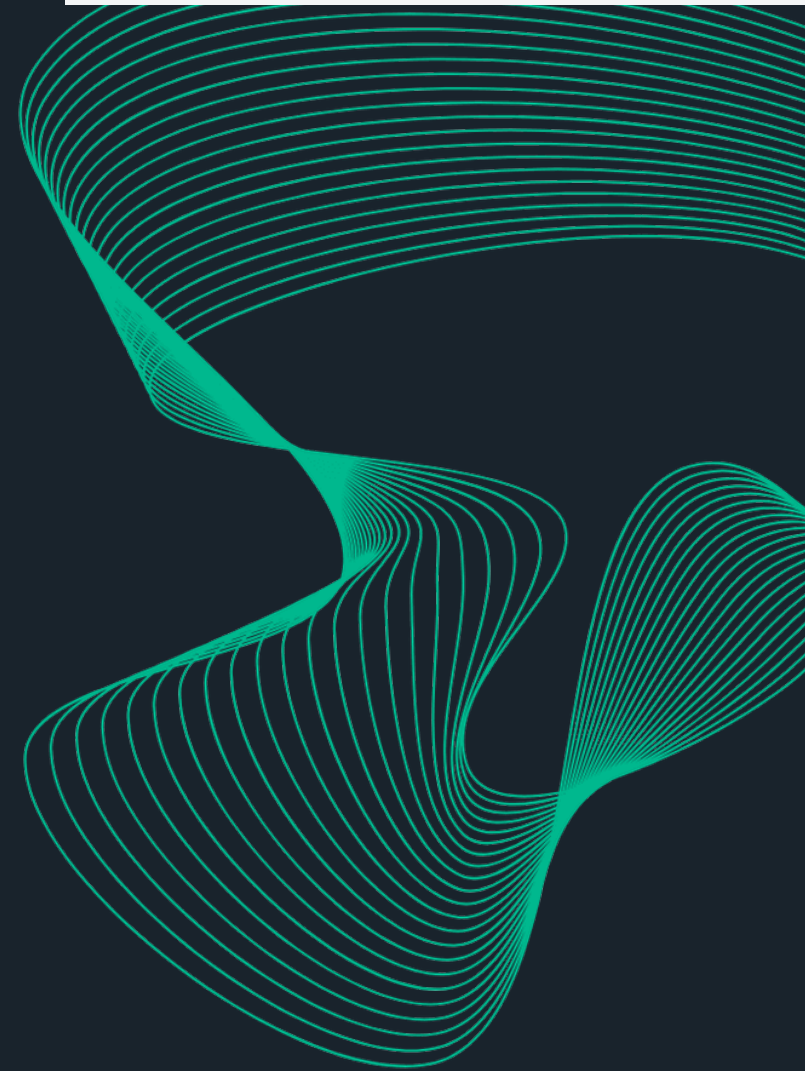
▶ [Click link to listen to audio](#)

AI Governance & Responsible AI

The European Union has enacted comprehensive legislation (the **EU AI Act**) that classifies AI systems according to the risks they pose, prohibiting those deemed unacceptable (e.g., social scoring) and regulating others that are high-risk. Developers of AI systems (called “providers”) have the bulk of the responsibilities, but deployers (called “users”) of such systems have some obligations as well. The Act also regulates developers of large AI models, requiring such entities to comply with European copyright directives and conduct certain evaluations, testing, and incident reporting.

The EU AI Act has extraterritorial reach, capturing AI systems whose outputs are used in the EU. The Act operates in tandem with the EU General Data Protection Regulation, with the latter controlling when there is a conflict between the two. In addition to the EU, other jurisdictions are developing regulatory frameworks, including Singapore, South Korea, China, and Brazil.

Finally, multiple organizations—including non-governmental organizations, government agencies, and companies—have developed principles for “responsible AI.” While these principles may differ in their specifics depending on the particular framework, they are generally designed to promote the human-centric design of AI systems; fairness and elimination of algorithmic bias; accountability and oversight; transparency; privacy and security; and safety and reliability.



Data Privacy

Processing personal data to train, fine-tune, or use AI systems can trigger state, federal, and international data privacy laws.

Personal data is broadly defined as information that is linked or linkable, directly or indirectly, to an identified or identifiable natural person, and it can include everything from images and audio recordings to IP addresses and device identifiers. Developers of AI systems use massive amounts of data to train and fine-tune their systems, and deployers use data to fine-tune models to suit their particular needs. To determine their obligations under privacy laws, both must know whether data they use includes personal data, the source of such data, the legal authority to collect and use the data, and what individuals were told regarding how their data would be used.

Developers and deployers of AI systems may collect data directly from consumers, license data from third parties, or scrape data from websites and other online platforms. Each of these sources present different legal implications under privacy laws and, making things more complicated, privacy

laws differ regarding what data they allow entities to collect from consumers, how they may collect such data, for what purposes they may use the data, and under what circumstances they may disclose the data.

And while state privacy laws exempt “publicly available information” from the definition of “personal data,” they define “publicly available information” differently, making it difficult to know what obligations attach to data that was scraped from the internet for AI model training.

In addition, privacy laws restrict the processing of personal data for “profiling,” which state privacy laws generally define as any kind of automated processing of personal data—such as by predictive or generative AI systems—to evaluate, analyze, or predict personal characteristics, such as health, preferences, interests, or behavior.

Data Security

The proliferation of AI systems has exacerbated existing cybersecurity threats and introduced new ones. Hackers can use AI systems to launch sophisticated cyberattacks (such as deepfake-based social engineering) at a greater scale. AI systems also have become the targets of both familiar and novel cyber threats. Attackers looking to steal troves of sensitive data may target AI systems and their foundation models and may launch AI-specific attacks, such as by injecting manipulated or malicious data into AI models to produce incorrect or harmful outputs.

There are no cybersecurity laws specific to AI systems in the United States. Companies must comply with broadly applicable cybersecurity requirements at both the state and federal levels in securing their AI systems—and in defending against cybersecurity risks posed by AI. At a high level, companies must implement and maintain “reasonable” and “appropriate” cybersecurity controls for AI and other systems. Some state laws have more detailed requirements, such as implementing access controls, system monitoring, and encryption. If an organization's AI system is affected by a breach of personal data or certain significant cybersecurity incidents, they may be required to disclose that incident under various state and federal laws.

In the EU, broadly applicable and comprehensive cybersecurity laws, including the Network and Information Security Directive (NIS2), govern the security of organizations’ networks and information systems—including AI tools. The EU AI Act has requirements related to cybersecurity of “high-risk” AI systems, but they are general. For example, those systems must prevent and respond to malicious attempts to manipulate training data or AI models. In both the U.S. and EU, highly regulated companies in industries like financial services and healthcare face additional, prescriptive requirements.

Numerous authorities have published guidance and risk management frameworks to help companies manage cybersecurity risks associated with the development, deployment, and use of AI systems, including the Federal Bureau of Investigation (FBI), the NIST, the National Security Agency (NSA), the Cybersecurity & Infrastructure Security Agency (CISA), and authorities from Canada, the U.K., Australia, and New Zealand.

Copyright

AI creates significant copyright compliance challenges, particularly related to how models are trained and outputs are used. Training large models often involves ingesting vast datasets that may include copyrighted text, music, images, or code. Without appropriate licenses or reliance on recognized exceptions, companies can face infringement claims. Courts and regulators are still grappling with how to apply fair use and other legal doctrines in the context of AI; initial guidance and decisions have already produced mixed results, and over 40 cases are pending in U.S. courts. All of this has created uncertainty regarding the use of copyrighted works for model training. And the high number of works necessary to train foundation models—as well as the statutory damages scheme—and potential for attorneys' fee awards under copyright law may make any finding of liability tremendously costly.

AI-generated outputs raise risks.

Outputs that either reproduce or are substantially similar to copyrighted works used in training could also lead to liability for the model deployer. And even where the output is novel, it is not clear whether the output is protected under the various intellectual property law frameworks (for example, the U.S. Copyright Office's position is that works created by AI systems are generally not eligible for copyright protection, but this issue is less settled in the U.K.).

Antitrust

The rapid adoption of AI raises new antitrust compliance risks. One such concern relates to data access and exclusivity: firms with privileged datasets may foreclose competitors, while data-sharing arrangements risk being construed as collusion.

Market power issues can arise when dominant players leverage control over infrastructure, models, or platforms to either tie services or to block rivals. AI may also heighten collusion risks.

For example, agreements among competitors to contribute competitively sensitive data to algorithmic pricing tools may have the unintended consequence of aligning prices across competitors, and training on sensitive data could facilitate unlawful information exchanges.

Other potential risks include standard-setting collaborations, which should be designed to foster interoperability without excluding rivals, and open-source initiatives should be structured to avoid anti-competitive licensing restrictions. There also is a risk of self-preferencing by AI-enabled platforms, such as prioritizing affiliated services in recommendations or rankings.

Confidentiality

Ensuring the confidentiality of an organization's sensitive and proprietary information should be a central pillar of any AI governance program. The consequences of a misstep can be severe and can include regulatory penalties and contractual breaches, as well as reputational harm and loss of competitive advantage.

Failure to properly implement confidentiality processes can lead to unauthorized disclosure of sensitive information in unexpected ways.

For example, if trade secrets are used as inputs for an AI system, and those inputs are incorporated into the AI system's training data without appropriate safeguards, the system could later generate outputs to third parties that inadvertently reveal or replicate the trade secret material. In this scenario, the organization's ability to protect its key intellectual property could be jeopardized.

Government Contractor-Specific Considerations

Organizations that sell, or hope to sell, AI solutions to the federal government, (and those that are subcontractors on federal government contracts) have additional compliance considerations. These include obligations in the Federal Acquisition Regulations and Defense FAR Supplement (which impose strict obligations related to conflicts of interest, procurement integrity, and protection of controlled unclassified information) and specific cybersecurity requirements in federal programs for cloud service providers and defense contractors.

Many agencies, including the Departments of Defense and Homeland Security, also have their own AI standards. Other issues that may arise include requirements to conduct Privacy Impact Assessments (to demonstrate that AI systems respect civil rights, particularly in a law enforcement or surveillance context), as well as export-control and national security-related requirements.

[▶ Click link to listen to audio](#)

How E&C Can Support Business in Getting Ahead of Risks

In order to meaningfully apply this patchwork of standards and best practices to mitigate AI risks, ethics and compliance professionals need to ask a number of questions in order to be effective in this role. While this list of questions must be tailored to a given organization and its risk profile and updated to reflect the evolving technology and regulatory framework, here are some key questions to consider:

How Is Your Business Developing or Deploying AI:

- Are you developing the AI system or deploying a third-party system?
 - If developing:
 - Are you training the AI system?
 - Do you have appropriate consent or another lawful basis for obtaining and using the data?
 - Is the data sufficiently representative?
 - Have you followed responsible AI principles?
 - Are appropriate AI system guardrails in place?
 - Have you complied with relevant AI regulations?
 - Who are the customers who will purchase, or on whose behalf you will use the AI solution?
- If the federal government, or if you will be a subcontractor on federal contracts, have you considered applicable regulations and flow-down requirements?
 - Are there any agency-specific rules?
- Do contracts with prospective customers have relevant requirements regarding system development, access, etc.?
- In what jurisdictions will you be using or selling the system?
- What laws or rules apply (AI-specific laws, other relevant laws such as GDPR, industry-specific regulations, e.g., financial regulations, HIPAA)?

Risk Assessment & Impact Analysis

- Is this a high-risk use case under applicable laws (i.e., the EU AI Act)?
 - If yes, do you have a plan to comply?
- Does the system make or inform consequential decisions about people?
 - If yes, is there meaningful human review?
 - Are processes sufficiently transparent for audit and legal defense purposes?
- Could the system impact human rights?
 - If so, what are foreseeable potential harms to individuals or communities, and have mitigation strategies been identified?

Model Integrity

- Has the system been tested for bias, robustness, and explainability?
- Are safeguards in place to prevent data leakage or adversarial manipulation?

Security and Technical Controls

- Have you assessed cybersecurity threats such as model poisoning or prompt injection?
- Is there a plan for continuous monitoring and incident response?

Employee Training

- Have employees been trained on the risks and limitations of the system?

Key Components of an AI Compliance Program

Once you understand how the business is developing or deploying the AI system and the related risks and issues, you will be well-positioned to establish an appropriately tailored and risk-based approach. Many, if not all, of the AI-related risks you identify can be addressed through enhancements to your existing compliance programs (e.g., data privacy, confidentiality, information security, etc.) Below are AI-specific compliance program components to consider:

Governance and Oversight

Policies and Standards

Risk Assessment and
Impact Analysis

Processes and Controls

Training

Monitoring, Auditing,
and Testing

Reporting and
Incident Management

Documentation

Governance and Oversight

- Establish a cross-functional AI committee with business, technical, security, legal, and compliance representation.
- Assign accountability at senior leadership and board levels.
- Develop a governance framework (leveraging the questions above) to apply to new AI products.
 - Frameworks should cover development, deployment, contracting, and use.
 - Define approval requirements and escalation protocols.
 - Governance should apply to systems built internally and those procured from third parties.
 - Gather input from technical and business stakeholders, not just legal and compliance.

Policies and Standards

- Wherever possible, incorporate AI risks and compliance requirements into existing policies and programs.
- Adopt a responsible AI policy covering issues such as bias testing, data quality, explainability, vendor obligations, and prohibited use cases.

Risk Assessment and Impact Analysis

- Implement a robust risk assessment process for high-risk use cases (e.g., employment, credit, finance, healthcare, surveillance).
- Develop plans to mitigate identified risks.

Processes and Controls

- Implement formal data classification protocols (e.g., public, internal, restricted, confidential) that clearly identify what data may and may not be used in AI workflows.
- Adopt appropriate technical and operational controls and safeguards to protect data and AI systems, including strong access controls.
- Implement a process for contracts to incorporate explicit confidentiality obligations into AI vendor agreements—for example, a prohibition against training AI models on a company's data without express written permission.



Training

- Train employees on acceptable use of AI and approved AI tools, as well as the risks, limitations, and responsibilities of AI use.
- Educate employees about the risks of shadow AI use.
- Ensure that safe, approved AI tools are available to prevent the use of shadow AI.

Monitoring, Auditing, and Testing

- Continuously monitor AI systems for drift, bias, and unintended outcomes.
- Use red teams to simulate adversarial misuse.
- Conduct threat monitoring for AI-specific cyber risks.
- Conduct internal and external audits when appropriate.
- Consider technical monitoring controls and detection mechanisms for shadow use, for example:
 - Network and endpoint monitoring
 - Data loss prevention tools
 - Usage analytics to identify large uploads to generative AI endpoints.

Reporting and Incident Management

- Publicize channels to report concerns.
- Ensure issues are appropriately investigated and addressed.
- Adopt a plan for how detected shadow use will be handled.

Documentation

- Maintain comprehensive records around:
 - System design
 - Data sources
 - Risk assessments and testing results
 - Governance processes and decisions.

▶ [Click link to listen to audio](#)

Conclusion

While AI is capable of having tremendous positive impacts, it is nonetheless a significant governance challenge. By asking the right questions and embedding AI risk considerations into existing compliance architecture, companies can effectively manage risks without hindering innovation. The companies that will be most successful are those that view AI governance not only as a defensive exercise, but as an opportunity to build trust, demonstrate integrity, and be true leaders in a rapidly evolving world.



Davis Wright
Tremain LLP



Stacey Sprenkel

PARTNER

Head of Compliance,
Ethics, Risk & Governance
Head of ESG &
Sustainability



Filipp Kofman

PARTNER

Co-lead of AI Team



Nancy Libin

PARTNER

Co-chair, Technology,
Communications, and Privacy
and Security Practice



Michael T. Borgia

PARTNER

Head of Information Security
Practice



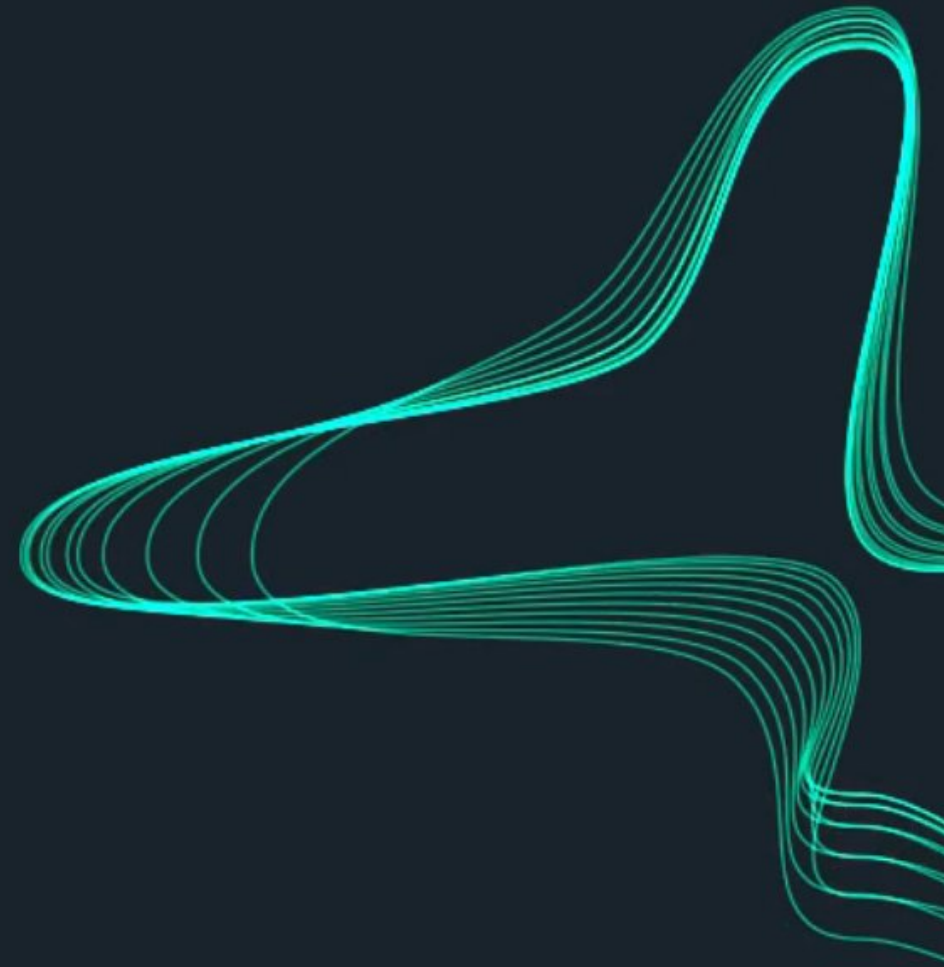
K.C. Halm

PARTNER

Co-lead of AI Team, Co-chair,
Technology, Communications,
and Privacy and Security
Practice

AI in Practice

Decision Rights, Controls, and Vendor
Coverage



Where Compliance Programs are Today

According to Ethisphere's data from World's Most Ethical Companies®, at leading E&C programs:

77%

have significant input or a coordinating role in responsible AI use

15%

include an AI provision in their Third-Party Code of Conduct

47%

have at least one E&C team member with a responsible-AI background

57%

have trained the general employee population on AI

The data shows that E&C is at the AI table and building literacy, but downstream governance (vendors, partners, and tools procured by the business) lags, exactly in the places where risk concentrates.

WHAT THIS MEANS:

- Decision rights exist, but artifacts lag. You can steer AI decisions, but without standard artifacts (model inventory, risk reviews, vendor clauses), your influence won't scale.
- eLearning ≠ control. Employee training is good, but control points (approvals, Human-in-the-Loop [HITL] triggers, monitoring) are what regulators and boards will look for.
- The blind spot is with third parties. Most AI capabilities that organizations leverage will arrive via vendors, only 15% coverage in Codes of Conduct is a material exposure.

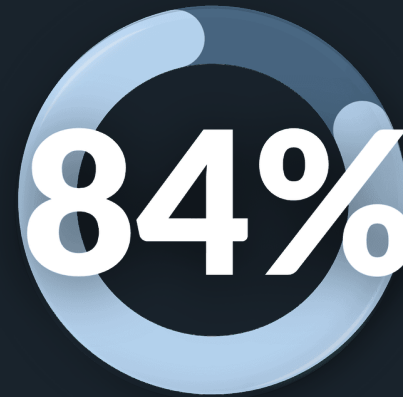
Where Programs are Today

Ownership without levers

E&C has the mandate: 84% of teams now own third-party risk management. But scale lags: only 58% routinely monitor high-risk vendors and just 14% have audited at least half of their third parties.

In other words, ownership is high while verification remains narrow.

Monitoring is concentrated at the top, audits are sparse, and AI clauses are rare. Closing the gap means moving from awareness to enforceable, AI-specific controls across the full vendor curve.



of E&C teams now own third-party risk management directly

58%



of companies routinely monitor higher-risk third parties

14%

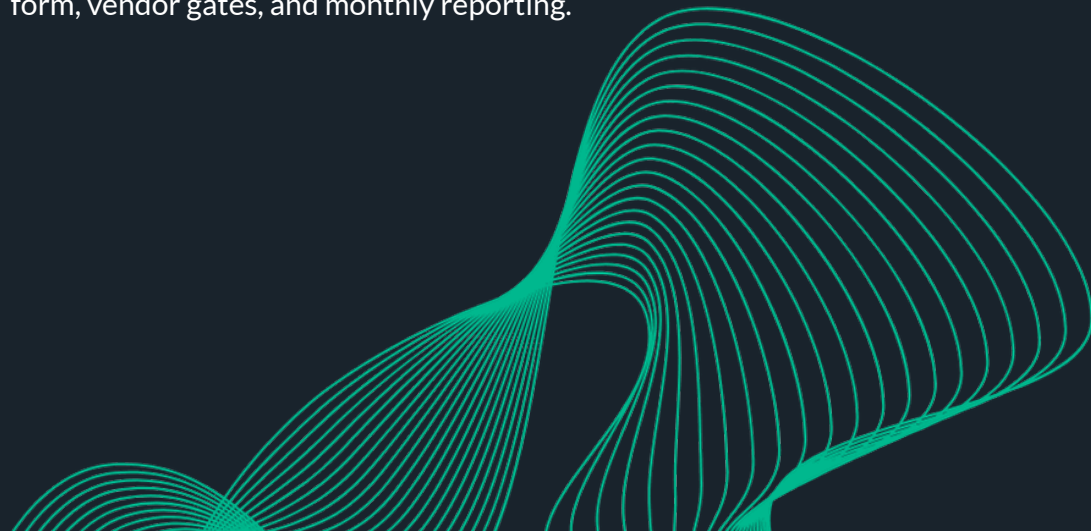


have audited at least half of their third parties

Governance to Enable Safe Scale

Decision Rights & Operating Model (RACI)

AI governance scales best when everyone knows their job. The following RACI assigns a single Accountable owner and clear Responsible, Consulted, and Informed roles for each governance activity. This allows E&C to move quickly without sacrificing control or questioning ownership. It's designed for cross-functional collaboration, including E&C, Legal, Risk, IT/Data, Procurement, and business leaders. And it should be mirrored in your intake form, vendor gates, and monthly reporting.



RACI USE CASES FOR AI IN E&C

Hotline intake summarization: who approves prompts/data, who reviews outputs before case creation

Policy & code translation: who validates accuracy/meaning, who maintains term glossaries

Third-party due-diligence scoring: who sets model features, who signs off on adverse-impact checks

Conflicts of interest patterning: who owns thresholds, who investigates automated flags

HR-adjacent scenarios (productivity/monitoring): who verifies legality and proportionality before deployment

Records retention & data governance for AI: who enforces deletion, access, and logging for model inputs/outputs

Executive reporting: who consolidates governance KPIs and incidents for leadership

ROLES & RESPONSIBILITES

R = Responsible; A = Accountable; C = Consulted; I = Informed

Activity/Decision	E&C	IT/Data	Risk	Legal	Business Owner	Procurement	AI Steering Group
Define AI policy and guardrails	A/R	C	C	C	I	I	C
Approve AI use cases (intake and triage)	A	R	C	C	R	I	C
Maintain model/use case inventory	A	R	C	C	I	I	I
Conduct risk review (per use case)	A	C	R	R	C	I	I
Define “human in the loop (HITL)” triggers	A	R	C	C	C	I	I
Ongoing monitoring & quality assurance (outputs, drift, bias)	A	R	C	C	C	I	I
AI-related incident management	A	R	C	C	I	I	I
General AI employee training	A/R	R	C	C	I	I	I
Role-based training (approvers/reviewers/creators)	A/R	R	C	C	I	I	I

ROLES & RESPONSIBILITES

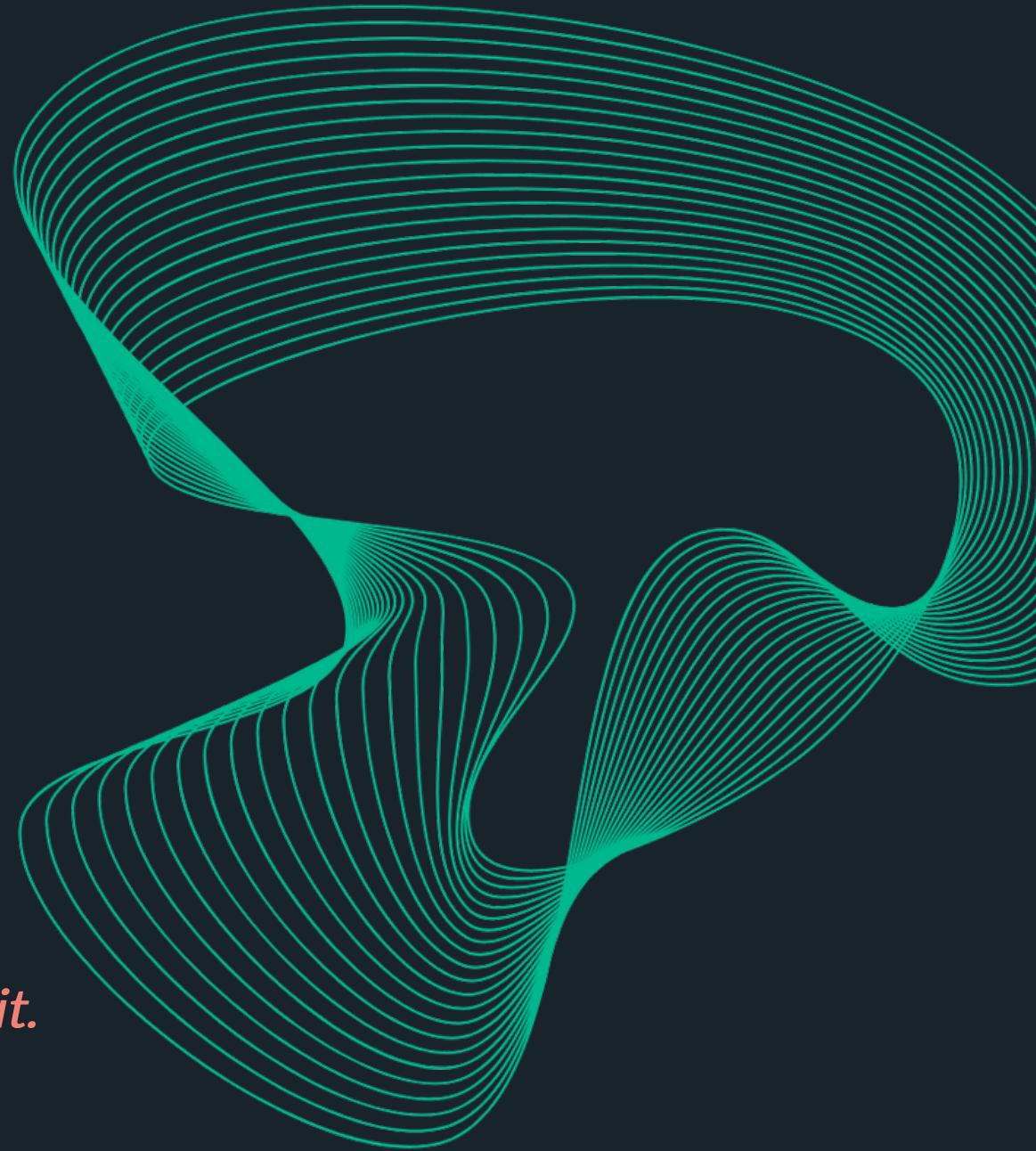
R = Responsible; A = Accountable; C = Consulted; I = Informed

Activity/Decision	E&C	IT/Data	Risk	Legal	Business Owner	Procurement	AI Steering Group
Vendor AI due diligence questionnaire	R	C	C	C	I	A/R	I
AI contract clauses and DPA updates	C	A/R	C	I	I	C	I
Vendor assurance review (testing, evidence)	A	C	C	C	I	R	I
Vendor performance monitoring for AI incidents	A	C	C	C	I	R	I
Data governance (privacy, retention, access)	C	C	C	A/R	I	I	I
Technical controls implementation (guardrails and tools)	C	C	C	A/R	C	I	I
Reporting to executives/Board (governance metrics)	A/R	C	C	C	I	I	I
Regulatory horizon scanning and alignment	C	A/R	C	I	I	I	I
Change management and communications	A/R	C	C	C	C	I	I
Budget and resources for AI in E&C	A/R	C	C	C	C	I	C

Baseline AI Governance Controls

Minimum controls create a documented trail of accountability: what you approved, why, the safeguards in place, and how outputs are checked. They're designed to be lightweight but testable and provide easy evidence to auditors, regulators, and your board. Implement them across the organization (including vendors) before expanding to advanced use cases.

Baseline ≠ Basic
Prove governance, don't just promise it.



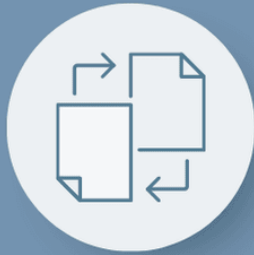
AI Policy + “Shadow AI” Addendum

Objective: To set clear guardrails for what’s allowed, what’s not, and which tools or processes are approved



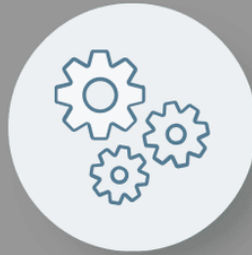
RACI

(A/R): E&C
(C): Legal, Risk,
IT/Data, AI Steering
Group
(I): Procurement,
Business Owners



ARTIFACTS

- AI policy
- Shadow-AI rules
- Approved tools list
- Prohibited uses
- Data handling rules
- Vendor expectations statement
- Exception request form



WORKFLOW

Draft → Review with
Legal/Risk/IT →
Approve in Steering
→ Publish on intranet
→ Announce via
comms & training →
Annual (or event-
driven) refresh



SLA

Exceptions decided
within X business
days.

Policy reviewed at
least annually or after
material regulatory
change.



DOCUMENTATION

- Final policy version
- Stakeholder sign-offs
- Employee acknowledgments
- Exception log



METRICS

- Policy acknowledgment rate
- Exception count and turnaround
- Shadow-AI incident trend

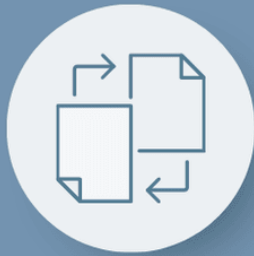
Model/Use-Case Inventory

Objective: To maintain a single source of truth for every AI use case and model touching E&C work



RACI

(A): E&C
(R): IT/Data
(maintains entries);
Business Owner
(keeps their row
current)
(C): Legal, Risk



ARTIFACTS

Spreadsheet (or similar)
with the following fields:

- Use-case name
- Business owner
- Purpose
- Model or tool & version
- Data sources
- Jurisdictions
- Approvals & dates
- HITL triggers
- Monitoring plan
- Go-live date
- Next re-review date
- Vendor & contract references



WORKFLOW

Intake auto-creates a
row → Reviewer verifies
completeness → Status
flips to “Approved” only
after Risk Review →
Regular (monthly,
quarterly, etc)
attestation by Business
Owner, with defined
regularly occurring E&C
spot-check



SLA

New or changed use
cases reflected in the
inventory within X
business days



DOCUMENTATION

- Inventory export
- Monthly attestations
- Change history



METRICS

- % of active use cases listed
- % with current approvals
- % of past due for re-review

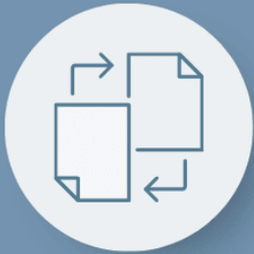
Risk Review

Objective: Conduct lightweight pre-go-live check that confirms legality, fairness, safety, and auditability



RACI

(A): E&C
(R): Legal and Risk
(perform review,
document decision)
(C): IT/Data, Business
Owner; Informed
(I): Procurement,
Steering



ARTIFACTS

- Short intake or review form
- Decision record (approve/conditional/deny)
- Re-review date
- HITL plan
- QA plan



WORKFLOW

Business Owner
submits form →
Legal/Risk review →
E&C issues decision
→ artifacts filed in
inventory →
conditions tracked to
closure



SLA

Decision within X
business days for
standard risk and Y
for high-risk.



DOCUMENTATION

- Completed form
- Signoffs
- Conditions log
- Re-review schedule



METRICS

- Cycle time from intake to decision
- Approval or conditional rates
- % with defined HITL
- QA plan

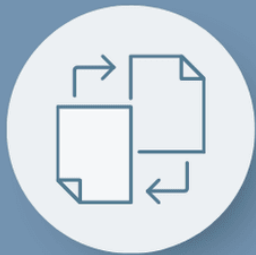
Human-in-the-Loop (HITL) Triggers

Objective: To ensure human review for high-impact or high-risk outputs before action



RACI

(A): E&C
(R): Business Owner (reviewers, SOP updates); IT/Data (technical routing)
(C): Legal, Risk



ARTIFACTS

- Activation ("trigger") list by scenario
 - (e.g., allegations with protected classes, employment actions, sanctions or TPRM escalations, high-risk geographies)
- Reviewer roster
- Escalation path
- Override authority
- Evidence checklist



WORKFLOW

Define activations → configure routing in tools → train reviewers → run UAT → go live → sample reviews monthly; tune thresholds



SLA

Reviewer response time (e.g., X business days standard; Y hours for priority cases).



DOCUMENTATION

- HITL decision logs
- Sampled outputs with reviewer notes
- Overrides



METRICS

- % of eligible items reviewed
- Exception or override rates
- Reviewer SLA adherence

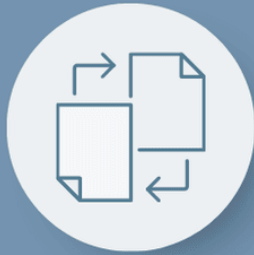
Monitoring & QA

Objective: To verify that AI outputs remain accurate, unbiased, and fit for purpose; and to catch drift early



RACI

(A): E&C
(R): IT/Data
(sampling/tests),
Business Owner
(operational QA)
(C): Legal, Risk



ARTIFACTS

- Sampling plan
 - (frequency, sample size)
- Acceptance criteria
- Test scripts or cases
- Exception log
- Incident playbook
- Change log for prompts, models, or data



WORKFLOW

Run sampling on schedule → log exceptions → root-cause with IT/Data and Business Owner → corrective actions → report in monthly governance pack



SLA

Exceptions triaged within X business days and material incidents escalated Y hours.



DOCUMENTATION

- QA results
- Exception tickets
- Incident postmortems
- Remediation proofs.



METRICS

- Exception rate over time
- Time to remediation
- Sampling coverage
- Material incidents per quarter

Vendor & Third-Party Governance

Third parties are where most AI enters the organization. They're also where risk is concentrated.

Yet, only 15% of programs currently include AI provisions in their Third-Party Code.

This section closes that gap with a simple, repeatable gate: ask the right questions before purchase, bind protections at contract, require assurance evidence before going live, and monitor continuously. Use the same controls for in-house and vendor tools so value is scaled without outsourcing risk.

- **Gate:** No AI-enabled vendor goes live without (a) an AI questionnaire, (b) contract clauses, or (c) assurance evidence
- **Questionnaire essentials:** Data provenance and retention, training data sources, testing/red-team results, bias and robustness methods, explainability, and rollback plan
- **Contract clauses to discuss with Legal and IT:** Permitted AI uses, data location and deletion, sub-processor disclosure, incident notification ≤ 72 h, right to audit or testing artifacts, prohibition on training on your data without explicit consent
- **Code of Conduct clause:** Add an AI clause in your Third-Party or Supplier Code of Conduct to:
 - Set a baseline expectation for every vendor upstream of contracts
 - Normalize disclosure and extend your governance standards across the supply chain
 - Give Procurement/Legal clear footing to require testing evidence, incident notifications, and data-use limits

Examples of Third-Party Codes with AI Provisions

[International Paper](#)

[Citigroup](#)

[BMS](#)

[Sanofi](#)

[Lockheed](#)

[UnitedHealth Group](#)

[Westinghouse Nuclear](#)

[Booz Allen](#)

[Thermo Fisher](#)

[Enbridge](#)

[Five9](#)

[TIAA](#)

Vendor AI Review Framework

Most AI will enter the organization through vendors, concentrating risk. The following review framework makes AI due diligence a standard gate by plugging into sourcing without slowing the business and giving E&C a consistent and defensible way to approve, track, and roll back vendor AI (if needed).



DOWNLOAD

[Full Framework Questionnaire](#)

GOVERNANCE

Does the vendor have an established AI Steering/Governance Committee and/or responsible roles assigned?

[VIEW ALL QUESTIONS](#)

BIAS & ROBUSTNESS

How are biases detected and remediated?

[VIEW ALL QUESTIONS](#)

DATA PROVENANCE/RETENTION

How is the data provided by the customer being used in the AI system (e.g., training, informing the service, providing direct responses or delivery)?

[VIEW ALL QUESTIONS](#)

EXPLAINABILITY

Does the AI system provide explanations for how it provided a particular output?

[VIEW ALL QUESTIONS](#)

TRAINING DATA SOURCES

What sources were used to train the AI model?

[VIEW ALL QUESTIONS](#)

ROLLBACK PROCESS/CHANGE MANAGEMENT

Does the vendor have a documented Change Management plan?

[VIEW ALL QUESTIONS](#)

Dashboard Metrics

GOVERNANCE ADOPTION:

- % of AI use cases in inventory
- % with completed risk review
- time from intake to approval

CONTROL EFFICACY:

- % of use cases with HITL
- exception rate
- incidents reported and resolved

THIRD-PARTY COVERAGE:

- % of strategic vendors with AI clauses
- % of vendors completing AI questionnaire
- remediation turnaround

CAPABILITY:

- % of employees trained
- number of SMEs with RAI background
- training refresher completion

AI Outcomes for E&C

When governance and regulations are aligned, E&C efforts can shift from building the systems to reaping its dividends. Introducing AI can have organization-wide impacts on lowering operating expenses, expanding capacity, and reducing residual risk.

AI OUTCOMES FOR E&C

The impact extends to E&C teams as well, creating an opportunity to make improvements where it matters – in time, cost, precision, and capacity.

Here's how integrating AI can make a difference:

BOTTOM LINE:

Integrating AI into ethics and compliance workflows can deliver organization-wide benefits while directly improving E&C team performance. The core value-add lies in greater efficiency, lower costs, improved precision, and stronger defensibility of decisions.

TIME ↓



Shorter cycle times across core workflows.

KPIs:

- ✓ Median time-to-decision
- ✓ Time-to-close

COST ↓



Lower unit costs for routine analysis and content work.

KPIs:

- ✓ Cost per task/output

FALSE ALARMS ↓ PRECISION ↑



Better signals, less noise.

KPIs:

- ✓ Precision (valid hits ÷ total flags)
- ✓ False-positive rate

SPEED TO DECISION ↑



Faster, better-informed approvals.

KPIs:

- ✓ Decision SLA adherence
- ✓ First-pass approval rate

COVERAGE ↑



More languages, records, and scenarios reviewed.

KPIs:

- ✓ Coverage rate (languages/regions/processes in scope)

AUDITABILITY ↑



Decisions you can defend.

KPIs:

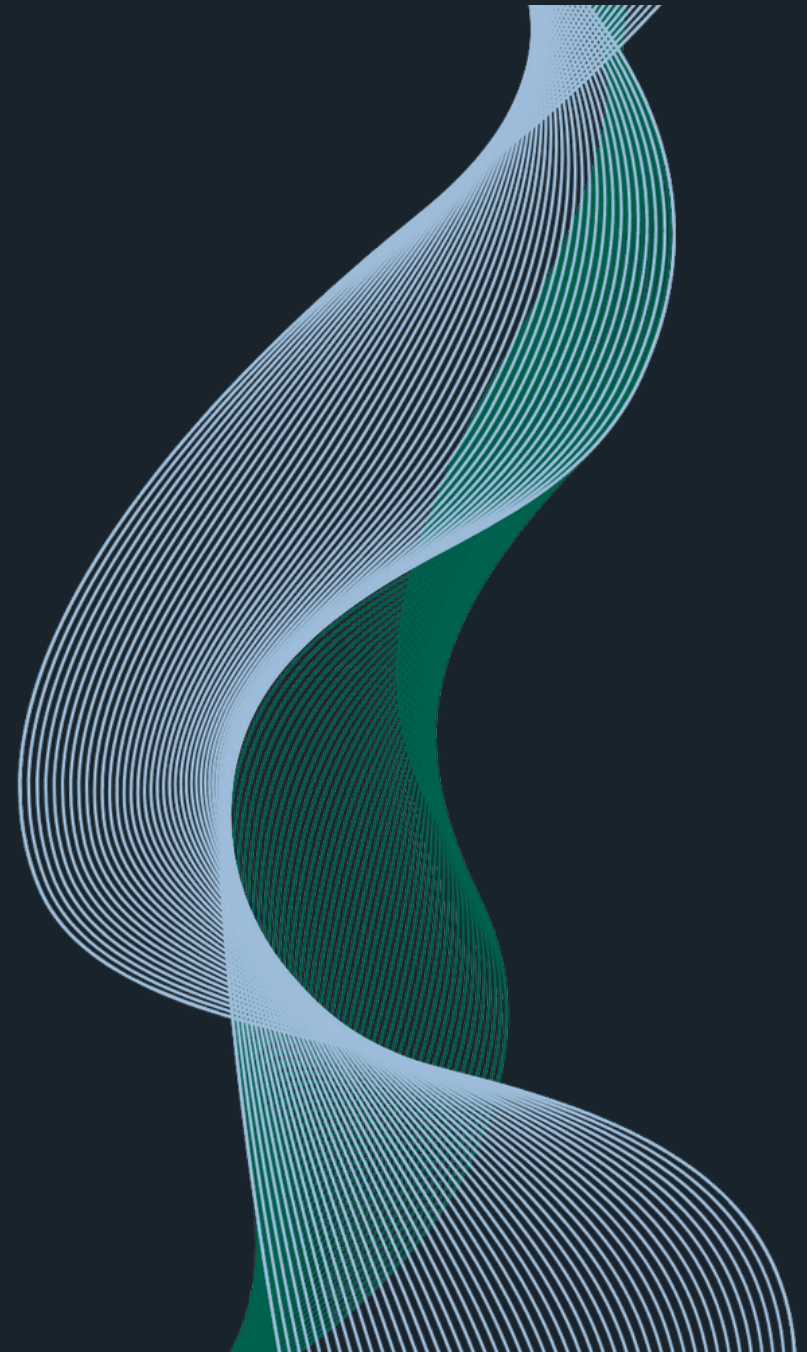
- ✓ % of approvals with evidence (form, HITL plan, logs)
- ✓ Version pinning in place

Emerging Use Cases

When it comes to program implementation, it can be hard to know where to start. Governance creates the framework for assigning ownership, making approvals, and analyzing evidence to ensure that workflows are consistent and repeatable.

This section outlines practical use cases, the controls that make them defensible, and concise steps to create measurable value. Choose the outcomes that matter most for your organization and design workflows based on risk and effort.

Build momentum with small wins and documented results.



Investigations Triage

WHAT IT IS

AI can translate, summarize, and classify new reports to give investigators structured facts rather than raw narratives. Routing and de-duplication help the right team act sooner.

WHY IT MATTERS

Less time is spent on intake and assignment, allowing more time to be dedicated to the investigation itself. Assigning related reports to the same investigator ensures consistency in the process, promoting organizational justice.

OUTCOMES

- Time ↓ (faster triage and shorter cycles)
- Cost ↓ (less manual sorting/translation)
- False alarms ↓ / Precision ↑ (better severity classification, linking of related reports)
- Speed to decision ↑ (clearer first pass for reviewers, better pattern visibility)
- Auditability ↑ (source + output logs retained)

GUARDRAILS

Do:

Ground triage in approved sources and keep the original reporter's submission attached to the case

Define the purpose and data scope during intake, noting ownership and review date

Require HITL review for severity classification, protected classes, and high-risk geographies

Test for accuracy, bias, and hallucinations before launch and after any material changes

Document prompts, versions, and decisions for defensibility

Don't:

Automate assignment or closure of high-impact cases without human review

Source facts from the open web or use unvetted data

Store/share personal information beyond what is allowed by your policy and retention rules

Deploy an AI workflow without a fallback or rollback plan

Train models using case data without explicit approval and data privacy controls

AI-Enabled Risk Signals in Text and Documents

WHAT IT IS

AI gives structure to unstructured content, including case narratives, emails to shared mailboxes, third-party documents, and disclosures, allowing patterns and hotspots to emerge earlier.

WHY IT MATTERS

E&C can avoid sifting through noise and instead focus their efforts on direct interventions with business units or vendors that need attention.

OUTCOMES

- False alarms ↓ / Precision ↑ (smarter patterning and grouping)
- Time ↓ (less sifting through noise)
- Speed to decision ↑ (earlier visibility of hotspots)
- Auditability ↑ (thresholds, prompts, and decisions logged)
- Coverage ↑ (more content sources brought into scope)

GUARDRAILS

Do:	Don't:
Define limits, data source scope, and thresholds with your Legal and Risk teams	"Turn it on everywhere" without defining the scope of data sources or assessing for privacy impacts when required
Require HITL for any action with substantial impact	Trigger sanctions or HR actions from red flags raised by AI without human review
Link or group related reports and explain and document the rationale for doing so	Treat unexplainable outputs as indisputable facts
Test for accuracy, bias, and hallucinations	Retain unnecessary, identifying, or sensitive personal data in the tool's knowledge or reference base
Document change control and define re-testing triggers	Engage in excessive or intrusive monitoring that goes beyond legitimate business needs or violates an employee's right to privacy

AI as Your Partner for Program Analytics

WHAT IT IS

Ask natural-language questions about program effectiveness or health regarding timelines, substantiation, disclosures, and more. Leverage AI to help you assemble analysis, visuals, and narratives for your presentations and reports.

WHY IT MATTERS

Stakeholders and leadership receive actionable insights instead of dashboards, and analysts have more time to delve deeper into the data.

OUTCOMES

- Time ↓ (fewer ad-hoc pulls and manual analysis)
- Speed to decision ↑ (board-ready answers faster)
- Cost ↓ (analyst time redirected to higher-value work)
- Coverage ↑ (consistent KPIs across more datasets)
- Auditability ↑ (traceable analyses from system of record)

GUARDRAILS

Do:

Aggregate and anonymize data, removing personally identifying information

Anchor answers to your system of record and retain clear documentation of the analysis and decision-making processes

Use consistent, approved, and documented KPI definitions

Require human review for conclusions

Document all prompts, datasets, and versions for retention and defensibility

Don't:

Present AI-generated insights as final without human validation or review

Export raw case data to external AI or business intelligence tools without data privacy controls

Redefine KPIs to fit a particular narrative

AI-Assisted Reporting

WHAT IT IS

Use secure voice or web intake assistants to guide reporters through a natural conversation or structured intake flow to capture all key details and generate a clear, structured case summary that is translated and transcribed for the investigator.

WHY IT MATTERS

Leveraging an automated intake process can lower friction for reporters, improve clarity and understanding for case handlers, and expand language access.

OUTCOMES

- Coverage ↑ (more languages, constant access)
- Time ↓ (structured, complete reports at intake)
- False alarms ↓ / Precision ↑ (fewer misroutes due to missing details)
- Speed to decision ↑ (ready-to-review case packets)
- Cost ↓ (reduced rework and intake effort)
- Auditability ↑ (transcripts and structured records)

GUARDRAILS

Do:

Provide clear consent, confidentiality options, and anonymity wherever allowed

Guide reporters to share all relevant information to understand the full picture and preserve the original submission in alignment with company retention policies

Restrict the tool's knowledge base to approved policies and prompts

Make the tool available in all languages used by employees, and ensure it can be accessed on different devices and is easy to use

Document interactions and versions and enable quality sampling

Don't:

Pull from external sources or the open web during the intake process

Collect unnecessary sensitive data beyond the legitimate business need or purpose

Ask leading questions that could steer, bias, or skew allegations

Automate assignment of severity categories without review

AI-Anonymization for Confidentiality at Scale

WHAT IT IS

Anonymizing reports by automating the removal or masking of personal identifiers in reports and attachments, in accordance with retention rules that apply in your jurisdiction.

WHY IT MATTERS

Automating the process lowers manual workload, strengthens privacy protections, and increases trust with reporters and other individuals named in cases.

OUTCOMES

- Time ↓ (no manual redaction)
- Cost ↓ (lower effort to meet retention/privacy rules)
- Auditability ↑ (documented masking actions and reversals)
- Residual risk ↓ (safer sharing via data minimization)

GUARDRAILS

Do:	Don't:
Define masking rules by case type and jurisdiction rules	Make masking irreversible when law or policy may require unmasking under authority
Test for over- and under-masking and document exceptions and fixes	Mask case handlers, approvers, or audit fields needed for accountability
Allow for controlled unmasking with role-based access	Rely on automation alone without human review or auditing
Apply retention, deletion, and documentation rules and processes consistently	Change retention or masking rules without notice or approval
Document version and rule changes	Store original submissions (whether masked or unmasked) in unsecured repositories

Policy Q&A Assistant

WHAT IT IS

Leveraging an AI assistant to answer policy questions in plain language, that is grounded in your approved language and includes links back to the actual policies or Code of Conduct.

WHY IT MATTERS

Employees would not need to find and read a policy to adhere to it. This allows employees to easily understand the expectations outlined in policies and the Code of Conduct and reduces questions to the E&C team.

OUTCOMES

- Speed to decision ↑ (instant, cited guidance)
- False alarms ↓ / Precision ↑ (fewer misinterpretation-driven escalations)
- Coverage ↑ (self-service access for all employees)
- Time ↓ (fewer repeat inquiries to E&C)
- Auditability ↑ (questions and citations are documented)

GUARDRAILS

Do:

Ground answers in the actual language of your policies and link to the applicable policy or Code of Conduct

Include disclaimers and escalate human review protocols for complex or high-risk queries

Document questions asked and use patterns to inform training needs

Ensure data security and adhere to your retention and data privacy and security policies

Version and test prompt templates and review for bias, accuracy, and clarity

Don't:

Allow the tool to invent policy guidance or offer legal advice beyond the scope of your policies

Pull from the open web or non-approved sources

Answer questions on topics outside the defined policy library

Make material changes to responses or models without following protocols or approval

Regulatory Change Intelligence

WHAT IT IS

Continuous monitoring of relevant laws and guidance with AI-assisted summaries that track changes and predict the impact those changes will have on E&C programs.

WHY IT MATTERS

E&C stays ahead of regulatory changes and can prioritize updates by impact rather than headline, staying proactive in identifying potential gaps in their programs.

OUTCOMES

- Time ↓ (faster impact mapping)
- Speed to decision ↑ (prioritized action lists)
- Coverage ↑ (more jurisdictions tracked consistently)
- Auditability ↑ (traceability from change → policy/control updates)
- Cost ↓ (targeted research vs. broad manual reviews)

GUARDRAILS

Do:

Curate authoritative, trusted sources and publish the list

Require human validation before actions

Map changes and regulatory updates to impact, owner, and due date

Document evidence and justification for changes, including citations, timestamps, and jurisdiction

Set a cadence and triggers for updates and document whether updates are adopted, monitored, or deferred

Don't:

Pull from unvetted or unreliable sources or the open web

Treat AI summaries as final legal positions

Release summaries of updates without assigning ownership or follow-through

Include excerpts without context or documentation, such as links, dates, or jurisdiction

Release summaries that include conflicting requirements without human review to determine how to resolve the conflicting laws

Compliance Knowledge Assistant

WHAT IT IS

A generative AI assistant for the E&C team that is grounded in your approved knowledge base of policies, procedures, controls, Code of Conduct, guidance, playbooks, training, prior board reports, and select legal memos. Use the assistant to answer questions, draft briefs, pull reference material, and assemble information for reporting.

WHY IT MATTERS

The assistant can reduce research time, improve consistency of guidance, preserve established practices and experiences, and accelerate onboarding without relying on the Legal team or senior SMEs to answer repeat questions.

OUTCOMES

- Speed to decision ↑ (quicker, cited answers for policy or process questions)
- Time ↓ (less ad-hoc SME research)
- Coverage ↑ (broader access to institutional knowledge)
- Auditability ↑ (documented questions and answers with sources)
- False alarms ↓ / Precision ↑ (fewer misinterpretations when guidance is cited)

GUARDRAILS

Do:

Ground responses in your approved knowledge base and require citations or links to resources

Enforce role-based access to sensitive documents to ensure data privacy

Document questions and answers and review them for gaps or policy update opportunities

Document evidence and justification for changes, including citations, timestamps, and jurisdiction

Test for bias, safety, and accuracy on any updates made to canned responses

Don't:

Pull from the open web or undocumented or unreliable sources

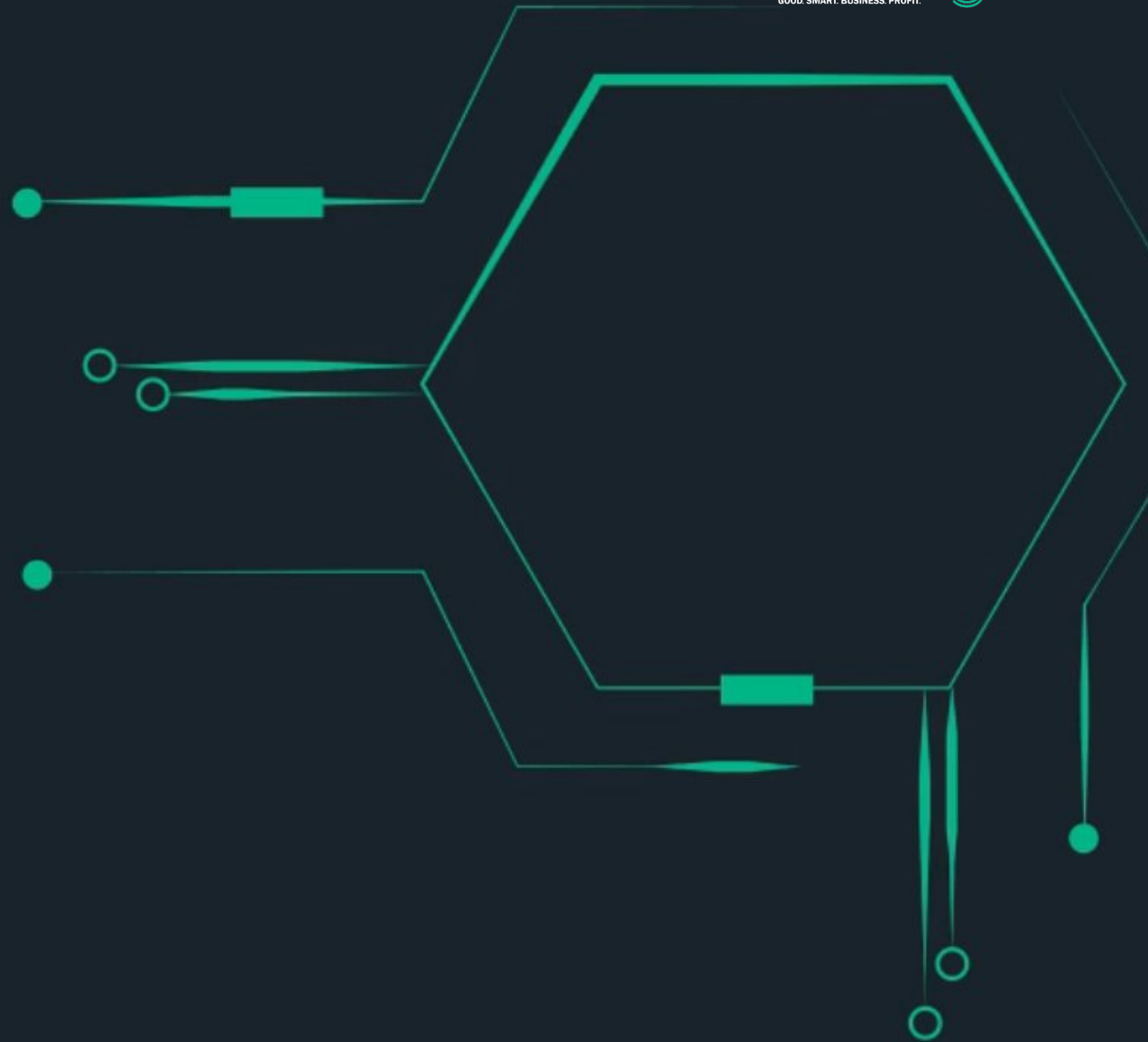
Allow privileged or draft materials to be accessed by individuals without a business need

Provide definitive legal advice on complex queries

Change responses or model answers without change control or notice

Using AI Inside the E&C Function

Insights from Business Ethics Leadership
Alliance (BELA) Organizations



Progress Over Perfection:

How Palo Alto Networks' E&C team is putting AI to work

Insights from Catherine Razzano, Global Ethics and Compliance Leader, Palo Alto Networks



Building internal momentum to adopt AI inside sensitive legal workflows

Culture helps. According to Razzano, Palo Alto's CEO set a clear expectation that every function would develop a real use case for AI. That mandate gave Legal and Compliance teams permission to experiment. She also challenges a common assumption: not everything handled by Legal and E&C is so sensitive that it precludes model-assisted workflows. With reasonable controls, many tasks are suitable for AI augmentation and will benefit from it.

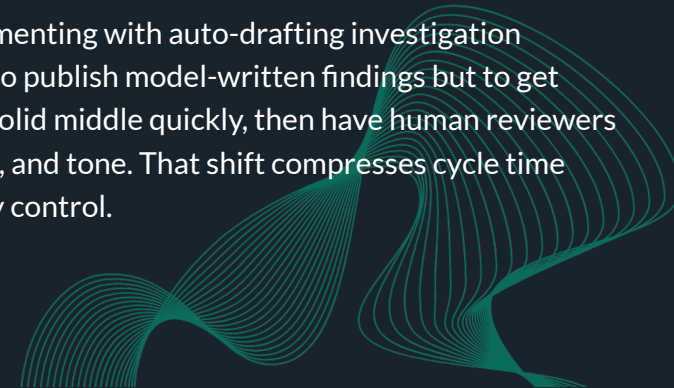
Transforming investigations

Razzano describes a pragmatic start that favors momentum over theory (or, as she puts it, "progress over perfection"). In particular, her team uses enterprise tools including Google's Gemini and NotebookLM to accelerate investigations. They load large, case-relevant data sets into a secure notebook environment—email, Slack, CRM records, purchase orders, and expenses, where privacy allows—then instruct the model on the allegation and ask it to surface the most relevant threads.

In effect, AI performs a first-level review that previously went to contractors or junior reviewers. Investigators then work from the model's prioritized outputs.

The payoff is time. Just as important, prompt discipline is sharpening investigative scoping. By forcing the team to state precisely what they are looking for, the tool narrows the inquiry to the facts that matter and the evidence most likely to advance the case. Privilege questions remain on the radar, but Razzano's preference is to move, learn, and manage risk rather than hold for perfection.

The team is also experimenting with auto-drafting investigation reports. The aim is not to publish model-written findings but to get from a blank page to a solid middle quickly, then have human reviewers tighten facts, reasoning, and tone. That shift compresses cycle time without relaxing quality control.



Where else is AI already delivering business value for Legal and E&C?

Razzano points to a simple but telling example. Outside counsel delivered a long memo on a regulatory issue. The business needed next-day, actionable steps. Instead of sending the memo back for a “plain English” rewrite, the team fed it into NotebookLM, asked the questions they would have asked counsel, and prompted the system for a checklist formatted the way their business partners prefer. Within 30 minutes, they had a usable plan that saved rework, outside spend, and a day of delay.

The result was closer to the business need than a revised memo would likely have been.

For Razzano, this illustrates three themes that resonate with boards and audit committees: efficiency, cost control, and better service to the business—all of which align directly with her team’s own operating principles. It is the kind of story that makes AI tangible for non-technical stakeholders.

How are you approaching employee engagement with AI tools?

Internally, the company is iterating a chatbot that does more than link to policies. Version 1 can answer “What is our gifts policy?” with a policy citation. Version 2 is designed to ask context questions (e.g., where a client dinner might be, whether the guest is a government customer or a commercial client), then guide the employee to the right rules while quietly capturing structured data about the interaction. The goal is a more conversational, consultative experience that raises engagement and generates operational signals E&C can use.

Guiding principles to operationalizing AI

Razzano points to a four-point outlook that provides a helpful and empowering context around how her team



Start and learn. Use AI where it can safely shrink effort and clarify focus.



Business orientation. Measure value by cycle time saved, dollars avoided, and whether outputs are truly usable by the business.



Human-in-the-loop. Reserve human judgment for findings, conclusions, and narrative.

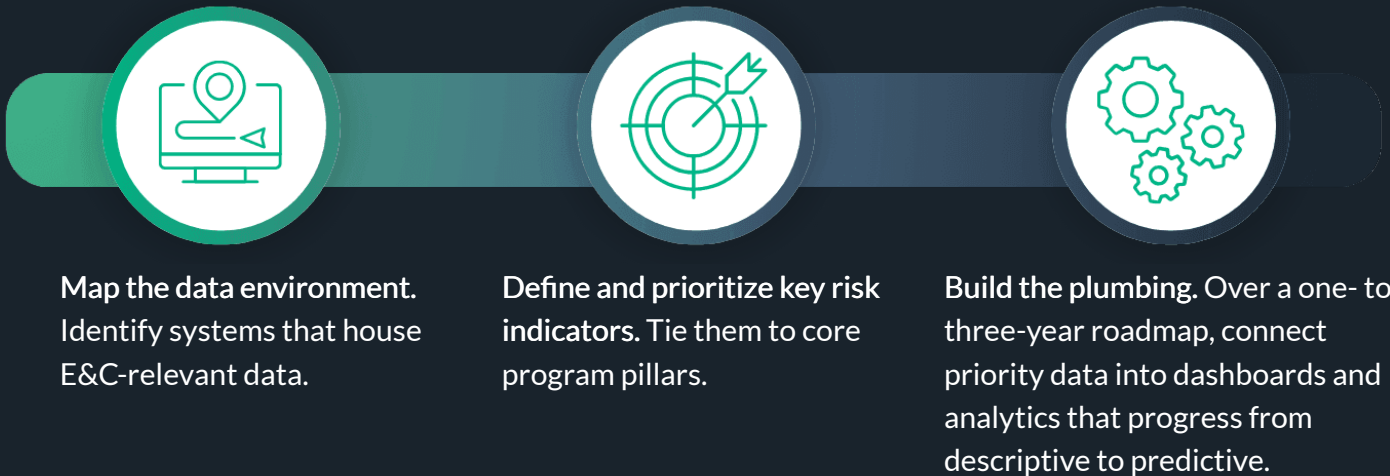


Data foundations. Treat isolated wins as steps toward an integrated, signals-driven program.

Where is this headed over the next 12 months?

Razzano expects continued expansion of AI-assisted investigations, wider use of model-drafted work products that humans finish, a more conversational internal chatbot, and steady progress on the data plumbing that enables risk indicators and dashboards. The predictive layer will take longer, but the groundwork is already underway.

For example, her team has already launched a data-utilization effort with three stages:



The north star is a data-driven program that flags risky transactions and patterns earlier, not a scattering of point solutions. Predictive capabilities sit at the far end of the roadmap, but the interim wins—visibility and faster decision support—matter now.

KEY TAKEAWAY

AI will not run the program.

It will let a lean E&C team cover more ground faster, with better alignment to the business. The mandate is to move with integrity and control while building the future state in parallel.

AI in E&C: From Literacy to Leverage

Insights from Samantha Vaughan, VP, Corporate Compliance Governance & Chief Privacy Officer, Verisk



Building literacy as a foundation

Vaughan emphasizes that effective, responsible, and ethical AI use starts with a strong literacy baseline. To that end, Verisk's privacy and compliance teams completed a four-hour AI fluency course, which covered not only AI ethics, but also:

- Large language model (LLM) mechanics
- Determining AI needs and the appropriate AI role (such as AI acting as a process assistant, consultant, or full delegate)
- Where AI excels and where it needs improvement
- Overreliance risks

The training was practical, requiring participants to experiment with prompts, compare outputs across models, and apply lessons to real or hypothetical compliance tasks. This investment created a shared vocabulary across the team and helped normalize AI as a daily tool rather than a novelty.

Modernizing compliance communications

At Verisk, an immediate AI application has been used to support reimagining their Code of Conduct. The existing version is solely text-based, similar to traditional policies. The new design in development emphasizes clarity, visual elements, and practical examples. Rather than drafting manually from scratch, Vaughan and her team are using AI as a virtual consultant at various points in the process. An internal, licensed LLM helps supplement best-practice examples, organize themes, and accelerate the drafting process.

It doesn't replace human authorship, but it reduces "blank-page work" and helps generate new angles for communicating core value expectations in practice.

Retrieval Augmented Generation (RAG) for training and enablement

One of the most valuable applications Vaughan identifies is Retrieval Augmented Generation (RAG). This approach allows compliance teams to point an LLM directly to internal documents (such as policies, codes of conduct, or guidance materials) and generate first drafts of training slides, blogs, or talking points. Within PowerPoint, for example, CoPilot can draw from corporate templates and designated policy documents to produce structured outlines and draft slides that practitioners can refine. While the AI output is not perfect, it offers an immediate and credible starting point that saves hours of manual drafting.

The key benefit is speed and efficiency, allowing compliance professionals to redirect time toward judgment, nuance, and stakeholder engagement.

Governance and the right tool for the task

Verisk's responsible and ethical AI governance framework recognizes that "AI" is not a single technology.



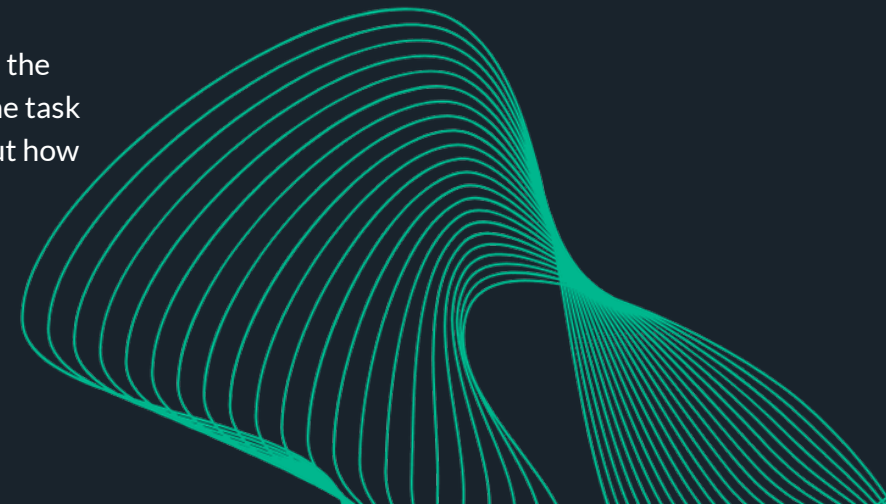
DOWNLOAD

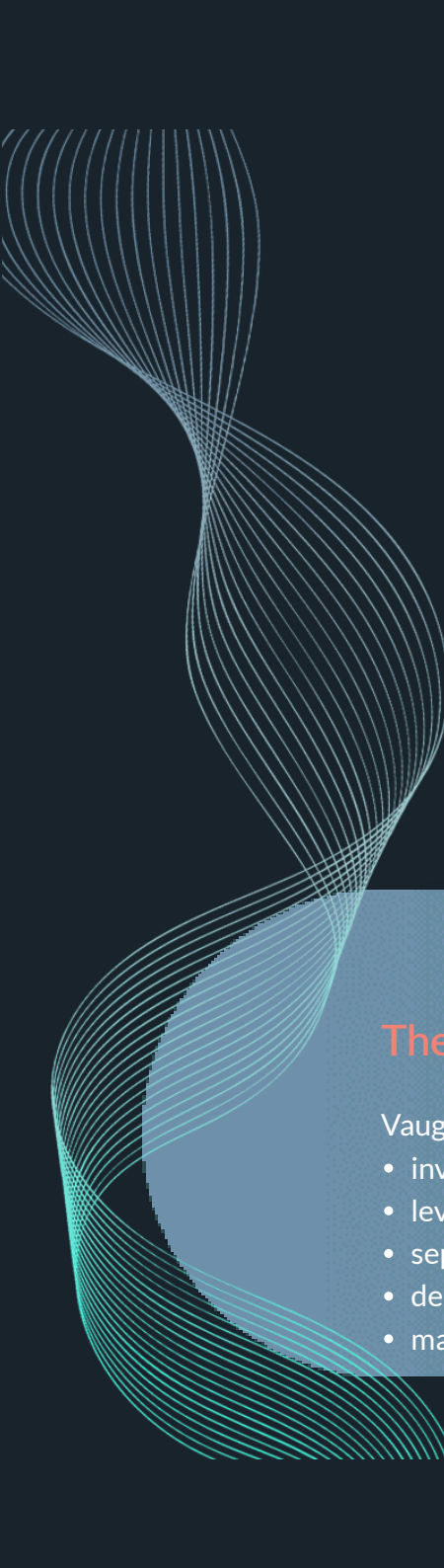
[AI Governance Framework](#)

Verisk maintains an ethical AI policy, plus distinct standards for generative AI and for algorithmic/predictive AI. Vaughan stresses that this distinction matters: predictive modeling, long used in insurance and other regulated industries, presents very different risks and oversight requirements than generative text. For compliance professionals, this underscores the importance of selecting the right tool for the task at hand and maintaining transparency about how models are being applied.

Vendor diligence as an ethical imperative

The growing reliance on external AI technology providers makes vendor diligence a core E&C responsibility. Vaughan cautions against treating vendor privacy assurances as sufficient overall diligence. Beyond whether personal data is used, compliance teams must insist on evidence of fairness testing, bias audits, and clear assessments of potential impacts on individuals, particularly in employment or HR contexts. Regulatory guidance (e.g., from the UK Information Commissioner's Office) provides useful benchmarks, but responsibility rests with the adopting organization. Outsourcing does not relieve accountability.





Evolving capabilities and rising risks

The rapid progress of AI capabilities continues to surprise even seasoned observers. Where early tools were rudimentary and error-prone, today's models can extract insights from even fragmentary materials, distilling meaning and generating useful questions or starting points. Yet the improvement also increases the risk of overconfidence: outputs may be polished and convincing, but still fundamentally flawed. Vaughan emphasizes the ongoing need for human-in-the-loop review to separate valuable insights from hidden errors or biases.

Upskilling as survival, not optional growth

Reflecting on her career, Vaughan notes that no prior technological wave—whether big data, the web, or modern productivity tools—demanded the same degree of mental transformation as AI. She believes professionals who fail to adopt AI will soon be outperformed by those who do, even at a basic level, and decisively outpaced by those who develop deeper mastery. The productivity advantage of tasks like generating first-draft presentations from policies is already visible. For E&C professionals, the lesson is clear: mastering AI is not an enhancement but a necessity for staying competitive.

The bigger picture for AI in E&C

Vaughan's perspective highlights a pragmatic blueprint:

- invest in literacy
- leverage RAG to accelerate real work
- separate governance across different AI modalities
- demand rigorous vendor diligence
- maintain human oversight at every step

Done thoughtfully, AI allows E&C teams to redirect time from drafting and formatting to designing stronger controls, sharper training, and faster, more informed outcomes.

Built-In, Not Bolt-On:

How Cargill's E&C Team Is Putting AI to Work

Insights from Samantha McMacken, Senior Director, Ethics and Compliance, Cargill



AI as a Thought Partner

Cargill's use of AI goes well beyond experimentation. The E&C team has built specialized agents, trained on policies and standards, that can draft investigation summary reports and support triage activities. These agents don't just generate formatted documents, they also prompt investigators with follow-up questions and potential angles that might have been overlooked.

The design goal is to make the agent behave more like a knowledgeable colleague than a passive tool, operating within strict guardrails and aligned with organizational principles.

Unlocking Data and Analytics

The technology has also become central to reporting and analytics. Preparing board materials and business dashboards once relied heavily on Power BI, requiring significant time and technical skill. Now, McMacken's team can upload data to GPT-based agents that cleanse, organize, and visualize the information according to Cargill's unique rules and calendars. Roughly 90% of visuals produced for the leadership reporting that McMacken manages are now generated through AI, with human review ensuring accuracy. This allows her to spot trends across longer time horizons, identify cultural red flags, and provide more nuanced insights than manual methods could deliver.

How to Train an Agent

Success depends on disciplined training. Rather than overwhelming agents with every possible resource, the team identifies each agent's core mission and loads only the most relevant materials. For example, the triage agent is trained on issue-type definitions, keyword maps, flagging criteria, and style guidance for summaries. Well-structured markdown instructions ensure the outputs are both accurate and usable.

Users are also encouraged to treat the agent like a colleague by questioning discrepancies, asking for explanations, and learning from the iterative exchange.

Keeping Humans in the Loop

While AI accelerates workflows, human judgment remains the final safeguard. Every output is reviewed, and responsibility ultimately rests with the investigator or analyst. This “human-in-the-loop” approach allows the team to maintain accountability while taking advantage of AI’s speed. The result is more communications, summaries, and reports completed in less time...all without sacrificing quality.



Faster Investigations

One area where this balance is particularly valuable is investigations. Intake and preparation tasks, such as adding summaries, tagging issue types, and formatting metadata, traditionally slowed down the process. With AI handling much of this routine work, investigators can focus on interviews, analysis, and decision-making. Upcoming integrations with transcription tools may further reduce cycle times by eliminating the need for manual note-taking, allowing investigators to manage more cases simultaneously.

Building Skills Through Curiosity

McMacken’s rapid adoption of AI was driven largely by experimentation. Before bringing the technology into her professional role, she tested different models in personal projects, studied external resources, and exchanged ideas with colleagues. She encourages peers to start small, using whatever tools their organizations approve, and to approach AI with curiosity. This hands-on exploration, she argues, is the fastest way to understand both the capabilities and the gaps that AI can address.

Elevating Visibility and Securing Support

AI adoption has also enhanced the visibility of the E&C function within Cargill. Early wins positioned the team as strong candidates for pilots of enterprise tools, including GPT Enterprise and Microsoft Copilot. Demonstrated value has since helped the team secure more licenses and ongoing support from leadership. Regular usage metrics and monthly surveys that quantify hours saved provide tangible evidence of impact, making budget discussions easier and reinforcing leadership confidence.



Elevating Visibility and Securing Support

AI adoption has also enhanced the visibility of the E&C function within Cargill. Early wins positioned the team as strong candidates for pilots of enterprise tools, including GPT Enterprise and Microsoft Copilot. Demonstrated value has since helped the team secure more licenses and ongoing support from leadership. Regular usage metrics and monthly surveys that quantify hours saved provide tangible evidence of impact, making budget discussions easier and reinforcing leadership confidence.

Distributed Impact

Beyond flagship initiatives, distributed adoption across teams has delivered meaningful gains. A colleague in derivatives, for example, applied GPT to automate a market surveillance task, saving more than 100 hours in a single quarter. These smaller, decentralized successes combine with larger efforts to create program-wide acceleration. The effect is not only efficiency but also broader buy-in, as employees see AI solving problems relevant to their own work.

Next Steps: Orchestration and Integration

Looking ahead, Cargill's E&C team is focusing on orchestration: linking multiple agents and automations for seamless end-to-end workflows. The vision is an integrated process where a bot can draw data directly from case management systems, triage it, draft reports, and return it to investigators without requiring manual transfers. Combining AI with existing automation platforms such as Power Automate will help ensure AI is "built-in, not bolt-on," as McMacken describes it, with humans in the loop as the final approver.

A Strategic Mindset

For McMacken, the philosophy is clear: AI should be embedded into systems and processes in ways that add lasting value, not patched on as an afterthought. That mindset, paired with disciplined training, measurable results, and a commitment to human oversight, is enabling Cargill's E&C team to keep pace with rising workloads, increase its strategic relevance, and pioneer a model for how compliance functions can responsibly leverage AI.



Epilogue

Data shows that E&C is at the AI table and building literacy, but **third-party governance lags**, creating material exposures exactly where risk is concentrated.



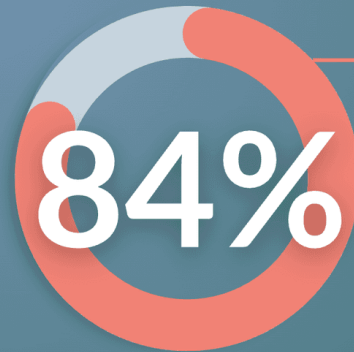
of E&C teams have significant input or a coordinating role in responsible AI use

57%

of E&C teams have trained the general employee population on AI

47%

have at least one E&C team member with a responsible-AI background



of E&C teams now own third-party risk management

Yet...

ONLY 15%

include an AI provision in their third-party code of conduct

ONLY 14%

have audited at least half of their third parties

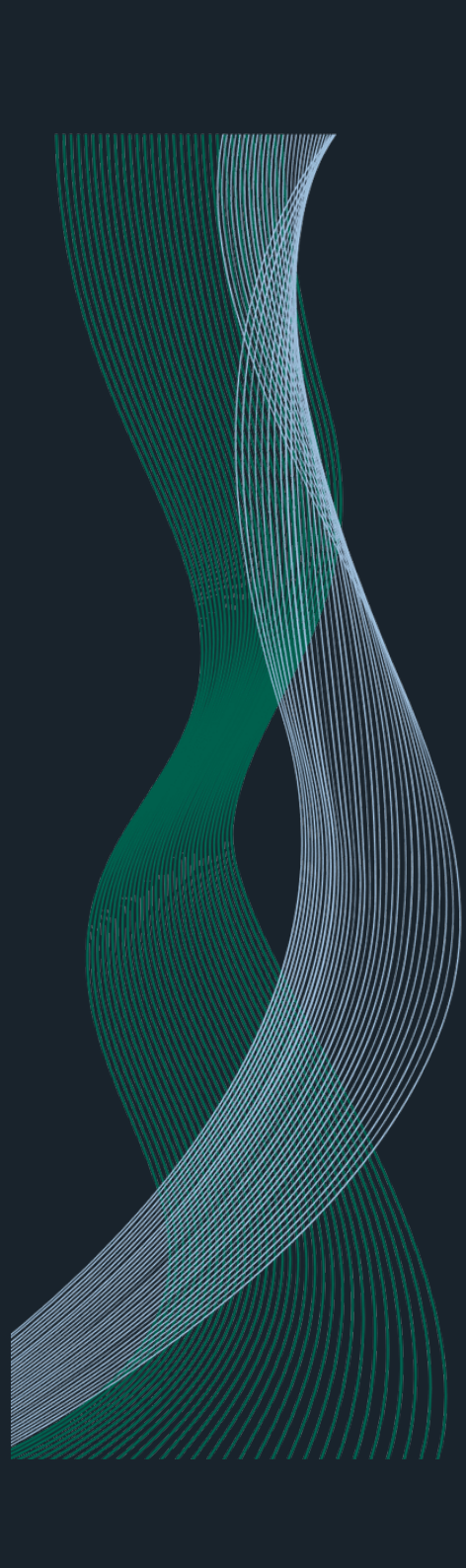
WHAT THIS MEANS

Decision rights exist but artifacts lag. You can steer AI decisions, but without standard artifacts (model inventory, risk reviews, vendor clauses) influence won't scale.

eLearning ≠ control. Employee training is good, but control points are what regulators and boards will look for.

The blind spot is with third parties. Most AI capabilities that organizations leverage will arrive via vendors, only 15% coverage in Codes of Conduct is a material exposure.

Ownership creates urgency. With 84% of E&C owning third-party risk management, there's no handoff excuse. The mandate is to standardize AI-specific controls and expand coverage in a way that scales.



Epilogue

As we bring this inaugural Artificial Intelligence Report to a close, we extend our thanks to the many ethics and compliance practitioners who shared their time, insights, and candor to help shape its contents.

Not everyone could participate. Several organizations told us that while they are actively exploring and implementing AI, the work is too sensitive and proprietary to be shared publicly. We respect and appreciate those boundaries. Even among the organizations that did go on the record, there was acknowledgment that this is an unusual moment: AI's rapid evolution has blurred the line between what is strategic for the enterprise and what is proprietary to compliance.

Compliance has long been a discipline where collaboration outweighs competition. Yet today, the transformative power of AI means that how an organization deploys it—whether for monitoring, investigations triage, or third-party risk—can itself confer competitive advantage. This will not always be the case. For now, we are in a rare window where AI adoption is experimental, entrepreneurial, and highly individualized. Every organization's use case looks different, and every success story is still being written.

We fully expect that when we return to this topic in six, 12, or 18 months, the landscape will look radically different. More case studies will exist, more best practices will be shareable, and the sophistication of AI integration in ethics and compliance will be far more advanced. What feels premature today will soon become standard.

One theme, however, has already emerged clearly:

Culture is central to AI adoption.

Across industries, leaders described how the strength of their culture—whether through a top-down mandate to experiment, or through an environment of psychological safety that encourages employees to test new tools—determines how well AI takes root. The same trust, accountability, and openness that sustain compliance programs also shape the success of AI initiatives.

Ethics and compliance has always been the steward of culture. Now, at this inflection point, culture is becoming the bridge that allows organizations to harness AI responsibly and effectively. The programs that are most successful today are those that leverage culture to unlock AI's

potential—advancing business integrity while positioning their organizations for the future.

This intersection of culture and technology is one of the most fascinating developments to emerge from this report. Going forward, it may well define the table stakes for what a truly effective ethics and compliance AI program looks like.

Thank you for reading. It has been a privilege to compile this report, and we look forward to revisiting this conversation soon.



Bill Coffin

Editor in Chief, Ethisphere



The Data and Organizations Behind This Report

Real-World Ethics & Compliance Data: For years, Ethisphere has built data sets that prove strong ethics is good business. Our program benchmarking data shows what excellence looks like for organizations seeking to develop best-in-class business integrity. And our Ethics Premium proves how much highly ethical organizations outperform their peers.

This report features insights from Ethisphere's proprietary data set, the 2025 World's Most Ethical Companies® part of the data intelligence BELA clients access to benchmark and strengthen their programs.

ETHISPHERE
GOOD. SMART. BUSINESS. PROFIT.®

Ethisphere is the trusted community for enabling global leaders to measure, mature, and unlock and demonstrate the power and value of their integrity programs. With 400+ global organizations, thousands of resources, and millions of ethical culture and E&C program data points, Ethisphere is re-defining how programs accelerate maturity.



SpeakUp is a leading provider of ethics and compliance technology, designed to simplify the way organizations manage risk and foster transparent workplace cultures. The platform enables employees to easily understand, report, and disclose ethics and compliance matters, while offering compliance officers a centralized system to manage daily workflows and act on critical issues efficiently.



Davis Wright Tremaine LLP (DWT) is a national, full-service firm representing businesses in the U.S. and globally. Known for a client-centric, industry-focused model, the firm blends multidisciplinary teams, technology, and deep sector insight to navigate complex litigation, transformative deals, and evolving regulation. BTI ranks DWT among firms best suited for high-value, novel work.

Thank you for reading

2025 Q3 - AI Report

This publication is informational and does not constitute legal, regulatory, or other professional advice. The content is current as of September 24, 2025 and may change without notice.

Ethisphere assumes no duty to update. Examples, benchmarks, and vendor references are illustrative only and not endorsements. © 2025 Ethisphere.