

Legal, regulatory, and insurance considerations when leveraging AI in investigations, RCA, and CAPA

Framing and operational guardrails for safety leadership
September 2026

[haven]

ReedSmith

Scope: U.S.-focused civil litigation, regulatory, and insurance coverage considerations
for AI-assisted investigations, RCA, and CAPA.
This document is general information and not legal, insurance, or coverage advice.

Contents

1. Executive summary	03
2. Legal exposure: Risk pathways and litigation mechanics	04
3. Regulatory exposure: OSHA and analogous regimes	06
4. Insurance and risk transfer considerations	07
5. Control framework: How to use AI and reduce avoidable exposure	09
6. Implementation checklist for safety leadership	10
7. System requirements for AI-enabled investigations, RCA, and CAPA platforms	11
8. Conclusion	14
Appendices	15

Authors

John Ellison – Reed Smith, Partner

Stephanie Gee – Reed Smith, Counsel

Joseph Hanna – Haven Safety AI, Co-Founder & CEO

1. Executive summary

Organizations are increasingly deploying artificial intelligence to support incident investigations, root cause analysis (RCA), and the development of corrective and preventive actions (CAPAs).

These tools have the potential to improve investigative rigor, accelerate pattern detection across incidents, and strengthen organizational learning. At the same time, legal, compliance, and risk management teams often raise an important question: Does the generation of more detailed analytical outputs and corrective action recommendations increase litigation, regulatory, or insurance exposure if those recommendations are not implemented?

From a legal and regulatory perspective, the use of AI does not introduce a fundamentally new liability theory. Courts and regulators have long relied on internal investigation reports, audits, engineering studies, and consultant recommendations to evaluate what hazards an organization recognized and how it responded. What AI changes is the scale, clarity, and accessibility of internal investigative records. AI-assisted systems can generate more structured analyses, alternative causal hypotheses, and recommended controls, which may become part of the discoverable record in litigation or regulatory inquiries.

Insurance adds a related but distinct layer of enterprise risk. Many commercial insurance programs were structured before AI became common in operational safety workflows, which can create ambiguity in policy terms, silent coverage gaps, and uncertainty over how claims involving AI-generated investigation content or CAPA recommendations will be handled. Emerging AI-specific exclusions in errors and omissions, professional liability, and commercial general liability policies make proactive policy review increasingly important.

The practical implication is not that organizations should avoid AI. Rather, organizations should implement AI within a governance framework that clearly distinguishes draft analytical outputs from approved conclusions; requires disciplined evaluation and disposition of recommended corrective actions; maintains appropriate controls over retention, privilege, and distribution of investigation materials; and includes a coordinated review of insurance coverage implications. When deployed with these safeguards, AI-enabled investigation systems can improve safety outcomes while maintaining a defensible and insurable record of how hazards were identified, evaluated, and addressed.

Scope note: This paper is U.S.-focused. It is written primarily through a civil litigation lens, with targeted sections on OSHA and analogous regulatory exposure and insurance coverage considerations. This paper also does not cover risks posed by AI models' bias, accuracy, or data integrity.

2. Legal exposure: Risk pathways and litigation mechanics

2.1 What legal counsel is concerned about

The concern most frequently raised by corporate counsel regarding AI-assisted investigations can be summarized in a simple phrase: *"your own system warned you."* In litigation, this concern generally manifests through arguments about notice, foreseeability, and the reasonableness of an organization's response to known hazards.

This dynamic is not unique to AI. Plaintiffs and regulators routinely rely on internal audits, consultant reports, engineering reviews, and traditional investigation files to construct similar narratives. AI systems do not create a new duty of care; however, they can significantly increase the volume and clarity of internal analytical records, making it easier for opposing parties to argue what the organization knew and what corrective options were available.

For safety leadership, the legal risk therefore arises less from the existence of AI-generated analyses and more from the absence of a structured process for evaluating and resolving those analyses. Systems that generate recommendations without clear ownership, prioritization, and documented disposition can create the appearance that hazards were identified but not addressed. Conversely, organizations that maintain disciplined review gates, clear documentation of decision rationale, and effective CAPA tracking are often better positioned to demonstrate that risks were evaluated and managed responsibly.

2.2 Discovery reality: Inadmissible does not mean undiscoverable

In federal civil litigation, discovery can reach nonprivileged matter that is relevant and proportional. The rules make clear that information need not be admissible at trial to be discoverable.¹

Operational implication: Even if the organization expects to challenge admissibility later, AI draft outputs can still be requested, reviewed, and used to shape deposition strategy if those draft outputs are nonprivileged and relevant.

2.3 AI artifacts are electronically stored information

AI-assisted investigation tools typically generate electronically stored information (ESI): Examples of such information include draft narratives, alternative causal hypotheses, recommendation sets, transcripts, comments, and audit logs. Requests for production commonly include ESI, and the mechanics of production are governed by Rule 34.²

Practical implication: Safety leadership should assume AI systems expand the set of potentially producible records unless privilege applies or retention rules control what is kept.

2.4 Internal rules and recommendations: Not usually the duty, but still relevant

Courts often resist turning internal policies into the legal standard of care. That limits the "you violated your own policy, so you were negligent" argument. At the same time, internal policies and recommendations can still be litigated as relevant evidence of knowledge, feasibility, and reasonableness, depending on jurisdiction and evidentiary rulings.⁹

Implication for AI: An AI recommendation generally does not become the legal standard of care simply because it exists. The more realistic risk is evidentiary – AI can make it easier to argue what was known and what controls were contemplated.

2.5 Subsequent remedial measures: Implementing CAPA is often less risky than assumed

Federal Rule of Evidence 407 generally restricts evidence of subsequent remedial measures when offered to prove negligence or culpable conduct, although it allows such evidence for other purposes, such as feasibility (if controverted) or impeachment.⁵

Implication: Fear of documenting improvements can be counterproductive. A defensible posture is to evaluate risk, implement reasonable controls, and verify effectiveness, while maintaining disciplined documentation.

2.6 Adoption and attribution risk

Draft AI outputs are different from approved company conclusions. If an organization adopts a statement, opposing counsel may attempt to treat the content as a party-opponent statement.⁶ Clear review gates, labeling, and separation between drafts and final reports reduce ambiguity.

2.7 Privilege and work product: Strong tools, but process dependent

Attorney-client privilege and work product are the most reliable protections for counsel-directed investigations. Corporate internal investigation privilege principles are grounded in *Upjohn*.⁷ Courts have also recognized that investigations can remain privileged even when they serve compliance purposes – as long as obtaining legal advice is a significant purpose.⁸ Work product doctrine traces to *Hickman* and is reflected in modern litigation practice.¹¹

AI does not break privilege, but AI can make privilege harder to defend if draft outputs are broadly shared, exported into ordinary business channels, or stored without lane discipline.

2.8 Do not design around a self-critical analysis privilege

Many jurisdictions do not recognize a broad self-critical analysis privilege as a reliable shield for internal evaluations.¹⁰ If the organization wants protection, it should use a counsel-directed lane where privilege and work product arguments have a stronger foundation.

2.9 Preservation risk scales with AI output volume

When litigation is reasonably anticipated, organizations must preserve relevant ESI. Federal Rule of Civil Procedure 37(e) addresses consequences when ESI that should have been preserved is lost because reasonable steps were not taken to preserve it.³ AI increases ESI volume, thereby increasing the need for clear retention rules, legal hold capabilities, and auditability.



3. Regulatory exposure: OSHA and analogous regimes

While this paper is primarily a civil litigation analysis, safety leadership should also consider how AI-assisted investigations and CAPA systems can shape exposure in regulatory inspections, enforcement actions, and post-incident inquiries. The dynamic is similar to litigation: The content itself is not the problem; the uncontrolled record and lack of disposition discipline are.

3.1 OSHA hazard recognition and feasible abatement narratives

OSHA's General Duty Clause enforcement is often framed around whether a recognized hazard existed and whether there was a feasible and useful method to correct it.¹³ If an AI system repeatedly identifies a hazard pattern and suggests plausible controls, those artifacts can be used to support hazard recognition and feasibility themes.

This does not mean AI creates new regulatory obligations. It means AI can create a cleaner documentary record of what the organization knew and what options it considered.

3.2 Willful framing and knowledge escalation

OSHA describes a willful violation as one where the employer knowingly failed to comply with a legal requirement (purposeful disregard) or acted with plain indifference to employee safety.¹⁴ OSHA's Field Operations Manual (FOM) provides additional guidance on willful violations and encourages consultation with counsel in developing such citations.¹⁵

Repeated internal findings that are not dispositioned or repeated high-risk recommendations that remain unaddressed can create a narrative of indifference. AI can amplify this risk if it produces more explicit recommendations without a corresponding system for evaluation, prioritization, and documented resolution.

3.3 Regulations that already require investigation and resolution of recommendations

Some regulatory frameworks explicitly require incident investigation reports, recommendations, and a system to promptly address and resolve findings and recommendations with documented resolutions. For example, OSHA's Process Safety Management (PSM) standard includes incident investigation requirements and requires a system to address and resolve findings and recommendations, with documentation.¹⁶

EPA's Risk Management Program (RMP) incident investigation provisions include a similar requirement to promptly address and resolve findings and recommendations and to document resolutions and corrective actions.¹⁷

In these environments, disposition discipline is not only a litigation control. It is also a compliance expectation. AI can help by improving completeness and traceability, but it can also increase the recommendation backlog if the organization lacks the capacity and governance to resolve the findings and recommendations.

3.4 Other sector regulators and post-incident inquiries

Beyond OSHA and EPA, many industries are subject to sector regulators or incident investigation authorities. Examples of such industries include mining, rail, aviation, pipelines, and energy. These other oversight entities can request internal records during investigations, and corrective action closure and effectiveness are common focal points.

4. Insurance and risk transfer considerations

Safety professionals evaluating AI-enabled investigations and RCA and CAPA tools should understand how existing insurance programs may respond – or fail to respond – when AI-generated content contributes to a loss.

Many commercial insurance portfolios were structured before AI became common in operational safety workflows. The result is potential ambiguity in policy terms, silent coverage gaps, and exclusions that may leave organizations exposed when AI-related claims materialize.

Organizations should evaluate these issues with qualified brokers, coverage counsel, and risk advisors. This section is intended to identify insurance governance considerations, not to provide coverage advice.

4.1 “Silent AI” coverage

“Silent AI” coverage refers to the condition where a policy neither expressly covers nor expressly excludes AI-related risks. Many general liability and professional liability policies currently fall into this category.

Safety professionals may assume that claims arising from AI-assisted safety decisions are covered under existing terms, but insurers may contest coverage on the basis that AI-generated recommendations constitute a distinct risk that was never underwritten or priced into the policy. This ambiguity may lead to coverage denials or delays in resolving claims.

4.2 Errors and omissions and professional liability exclusions relevant to AI-generated content

Professional liability and errors and omissions policies typically respond to claims arising from negligent acts, errors, or omissions that occur in the performance of professional services.

Recently, the market has seen the introduction of AI-specific exclusions in the context of errors and omissions and professional liability policies aimed at excluding losses arising out of, or in any way involving, an insured’s actual or alleged use of generative artificial intelligence.

If AI-generated materials are used as part of investigations, RCA, or CAPA, safety professionals should review their policies to determine whether such an exclusion exists and, if it does, whether it can be negotiated out of or narrowed in the policy.

4.3 Commercial general liability and the bodily injury or property damage nexus

Safety professionals may rely on commercial general liability (CGL) policies to respond to claims for bodily injury or property damage arising from an occurrence.

When AI-generated content contributes to physical harm, for example, and an AI system recommends a control measure that proves inadequate and a worker is subsequently injured, the CGL policy may respond because the ultimate harm is bodily injury.

Like errors and omissions and professional liability policies, the market has seen an increase in AI-related exclusions. For example, the Insurance Services Office (ISO) has introduced a standard endorsement designed to exclude coverage for bodily injury, property damage, and personal and advertising injury arising out of generative artificial intelligence (GenAI) in commercial general liability policies.

4.4 Proactive policy evaluation can be key for ensuring coverage

AI risks are unique to each business, depending on how AI is being used and the governance measures adopted for the use of AI. Safety professionals who proactively conduct thorough audits identifying these risks may find themselves in a better position with underwriters who have keyed in on AI risks and who may conduct more diligence on policyholders’ use and governance of AI.

Along with business-specific AI audits, safety professionals should review their insurance programs holistically in order to identify potential coverage gaps resulting from systematic integration of AI.



5. Control framework: How to use AI and reduce avoidable exposure

5.1 Establish the core principle in the safety management system

A defensible policy position is simple: AI generates candidate hypotheses and candidate CAPAs; humans approve findings, conclusions, and final CAPA decisions; only approved outputs represent the organization's position. This reduces ambiguity about adoption and supports consistent governance.⁶

5.2 Implement a two-lane operating model

Lane A, operational learning: routine events where litigation is not reasonably anticipated. AI can accelerate evidence organization, timeline reconstruction, and CAPA ideation. Outputs are ordinary business records, so write and store them accordingly.

Lane B, counsel-directed investigations: severe incidents or credible litigation anticipation. Counsel directs purpose, scope, and communications; access is restricted; workspaces are segregated; distribution is minimized. This aligns with both privilege and work-product defensibility.^{7, 8, 11}

In Appendix 1, we propose an example of a Lane B trigger guide.

5.3 Separate artifacts into three classes

To reduce misinterpretation and to manage the footprint of drafts, separate records into three classes: (1) evidence record, (2) AI working outputs, and (3) final approved record. This separation also makes ESI production more intelligible if it occurs.²

5.4 CAPA triage and disposition discipline

For every material recommendation, require a disposition state: adopt; adopt with modification; reject with rationale; defer with a time-bound re-review; or mitigate differently with an alternative control documented. For PSM – or RMP-covered sites, this discipline aligns with explicit regulatory expectations to address and resolve findings and recommendations with documented resolutions.^{16, 17}

5.5 Review gates and explicit sign-off

Create explicit approval gates before content becomes official: investigator review, operational owner review, safety leadership sign-off for high-risk CAPA, and counsel review for Lane B matters. This reduces the risk that drafts are treated as corporate conclusions.⁶

5.6 Write as if it will be discoverable

Given the breadth of discovery, adopt writing norms that withstand external scrutiny: separate facts from hypotheses, label hypotheses as hypotheses, avoid speculative intent and blame language, and reserve legal conclusions for counsel.¹

5.7 Retention and legal holds must scale to AI ESI

Define retention periods for AI working outputs, implement legal hold capability, and maintain audit trails showing preservation steps. This reduces spoliation risk in high-stakes events.³

5.8 Control access and distribution

Use role-based access controls, least-privilege defaults, and restricted export permissions for counsel-directed matters. Many privilege disputes and narrative problems arise because of casual sharing.

5.9 Messaging discipline

Avoid describing AI as a guaranteed prevention backstop. The safest framing is that AI assists with evidence synthesis and proposes candidates, while humans decide, implement, and verify. This reduces reliance on style narratives.¹²

6. Implementation checklist for safety leadership

6.1 Governance and decision rights

- Name an accountable executive owner (enterprise EHS leader).
- Define Lane A versus Lane B triggers and escalation paths.
- Define who approves findings and CAPA decisions (named roles).
- Align with legal on counsel-directed protocol and segregation requirements.^{7,8}

6.2 Process standardization

- Standardize the investigation pack: evidence register, timeline, findings summary, recommendation disposition log, verification plan.
- Require disposition and ownership for every high-risk recommendation, including AI-generated candidates.
- For PSM – or RMP-covered sites, ensure the disposition log supports the requirement to address and resolve findings and recommendations with documented resolutions.^{16,17}

6.3 System requirements

- Section 7 below is dedicated to the important topic of what technical requirements should be met with any system that leverages AI in the investigation, RCA, and CAPA workflow.

6.4 Training and quality assurance

- Train investigators on fact versus hypothesis separation and AI draft labeling.
- Train leaders on disposition discipline and effectiveness verification expectations.
- Quarterly quality assurance audits on time to disposition, overdue high-risk items, verification quality, and repeat events tied to unverified controls.



7. System requirements for AI-enabled investigations, RCA, and CAPA platforms

Safety leaders should treat an AI-enabled investigation system as a **high-consequence system of record**. It will influence prevention outcomes, shape organizational learning, and, in many cases, become part of the evidence footprint in civil litigation and regulatory matters. The right system requirements are therefore not just usability features. They are **governance, defensibility, and operational control features**.

This section defines what safety leaders should expect from an AI-powered investigation platform that supports investigations, RCAs, and CAPA recommendations, while meeting the legal and regulatory exposure constraints discussed in this white paper.

The discussions focus on the features needed to support the legal and regulatory frameworks discussed in earlier sections. It does not cover the general requirements for AI adoption, such as guardrails against hallucinations, PII leakage, model stability, and model drift.

This section also discusses recommendations for safety leaders on how to integrate insurance risk management into the design and governance process of the AI-enabled investigation from the outset.

7.1 Investigation workflow capabilities

Optimally, the platform should support an end-to-end investigation life cycle without forcing teams to stitch together various systems outputs, email, spreadsheets, shared drives, and slide decks. Stitching, copying, and pasting, and having the investigation records distributed across multiple systems and file stores – each with its own AI features – will only complicate the task of tracing the process and protecting the integrity of the investigation.

The scope should cover case intake and scoping, evidence capture and registration, timeline reconstruction, interview and statement collection, causal analysis and RCA, and CAPA management.

7.2 Governance design for defensibility

AI-enabled investigation systems generate many drafts, hypotheses, and intermediate artifacts. Without structured governance, the result may be a large, confusing record that is more difficult to manage or defend. Safety leaders should therefore expect the platform to enforce a clear governance structure.

A key requirement is support for a two-lane workflow model. The system should allow investigations to be designated from the time of creation or upon escalation as either routine operational cases or counsel-directed cases at creation or escalation. Each lane should operate in a separate workspace with distinct access controls, templates, review gates, and, where policy allows, different retention and export rules.

The system should also enforce a clear separation between draft analytical material and final investigation conclusions. AI-generated hypotheses, draft CAPA options, and investigator working notes should remain in draft states until formally approved. Approval workflows, visible labeling of draft content, and version-controlled status states should ensure that only approved findings and CAPAs appear in final reports. To reduce accidental distribution, the platform should also prevent the bulk export of draft material by default.

7.3 Evidence integrity and chain-of-custody features

Investigation systems must preserve the integrity of evidence even when the primary objective is operational learning rather than litigation. Safety leaders should expect the platform to maintain clear proof for all evidence collected during an investigation.

The system should preserve the original files and associated metadata, such as time stamps, source devices, uploaders, and locations, where available. Evidence objects should be stored in a way that prevents unauthorized alteration, using integrity checks or hashing mechanisms to demonstrate that files have not been tampered with.

The platform should also maintain a chain-of-custody log showing who collected, accessed, transferred, or exported each piece of evidence. Where sensitive information, such as personal or medical data, is involved, the system should provide redaction tools that protect the original record while allowing the controlled sharing of redacted versions.

7.4 Retention, deletion, legal holds, and e-discovery readiness

AI-enabled investigation systems generate significant volumes of electronically stored information, so the platform must support clear retention and preservation controls. Safety leaders should expect configurable retention schedules by artifact type, including evidence files, draft AI outputs, approved investigation reports, audit logs, and CAPA records. These policies should be flexible enough to align with corporate record retention requirements across regions and incident categories.

The system must also support legal holds. When litigation is anticipated, authorized users should be able to place a case or a group of cases on hold, automatically suspending deletion of relevant artifacts. The system should record when the hold was applied, who initiated it, and which records are preserved.

7.5 Regulatory exposure readiness requirements

Investigation systems should also support consistent responses to OSHA and to other regulatory inquiries. The platform should provide a role-based reporting view that assembles the factual elements of an investigation, including the timeline, evidence register, investigation conclusions (if disclosed), CAPA decisions, and relevant procedures or training references.

The system should also help to demonstrate that hazards were addressed systematically. This includes clear records of hazard identification, CAPA prioritization, implementation tracking, and verification of corrective actions.

Finally, the platform should enforce consistent taxonomies across sites so the organization can demonstrate enterprise-level learning and trend management. Role-based access controls should also prevent privileged or counsel-directed materials from being included in routine regulatory reports.

7.6 Integration of insurance risk management into governance

Governance policies and procedures should include an incident response protocol that considers insurance implications in the event of an AI-related failure or claim. Thinking about these issues proactively on the front-end will not only protect the company by giving consideration to risks that might be avoided as part of the development and implementation process but will also protect the operations if any adverse events develop once the processes are up and running. This process includes establishing clear protocols for providing timely notice to insurers and the preparation of claims submissions. Failure to integrate insurance notice and claims preparation into the broader incident response framework can result in delayed reporting, incomplete submissions, and, ultimately, coverage disputes when the organization most needs its insurance program to respond.

AI governance policies and procedures are not only operational necessities – they are increasingly relevant to insurance placement and underwriting. Carriers evaluating an organization's AI risk profile will look for evidence of established governance frameworks, including documented protocols for human oversight of AI outputs, performance monitoring, and data quality controls. Safety professionals who can demonstrate established protocols before any incident occurs as well as mature governance and controls around their AI-enabled safety systems may be better positioned to obtain coverage, negotiate favorable terms, and manage premium costs.

AI risk profiles are not static. As safety professionals expand their use of AI, modify their models, integrate new data sources, or extend AI-generated recommendations into new operational areas, the underlying risk profile may shift in ways that outpace existing insurance coverage. Safety professionals should build routine policy audits into their AI governance systems, establishing regular time intervals for reevaluating whether their current insurance programs adequately address evolving AI exposure.



8. Conclusion

AI-assisted investigation and RCA systems do not, by themselves, create new legal or regulatory liabilities. Courts and regulators continue to focus on familiar questions: What hazards were recognized, what corrective measures were available, and whether the organization acted reasonably in response to the information it had. The introduction of AI primarily affects the **volume, structure, and traceability of investigation records**, which can make internal analyses more visible in litigation or regulatory review.

For safety leadership, the appropriate response is not to limit the use of AI but to implement it within a disciplined governance framework. Organizations that clearly distinguish draft analytical outputs from approved conclusions; require structured evaluation and disposition of corrective action recommendations; and maintain strong controls over access, retention, and privilege are better positioned to demonstrate that they responded thoughtfully to safety risks.

When implemented with these controls, AI-enabled investigation platforms can strengthen both prevention and defensibility. They can improve the thoroughness of investigations, enable earlier detection of recurring hazards, and support more consistent corrective action management across sites. In that sense, the same capabilities that raise governance considerations can also provide organizations with clearer evidence that safety risks are identified, evaluated, and addressed in a systematic and accountable manner.



John Ellison
Reed Smith
Partner



Stephanie Gee
Reed Smith
Counsel



Joseph Hanna
Haven Safety AI
Co-Founder & CEO

Appendices

Appendix 1. Lane B trigger guide (example)

Escalate to a counsel-directed lane when one or more of the following apply:

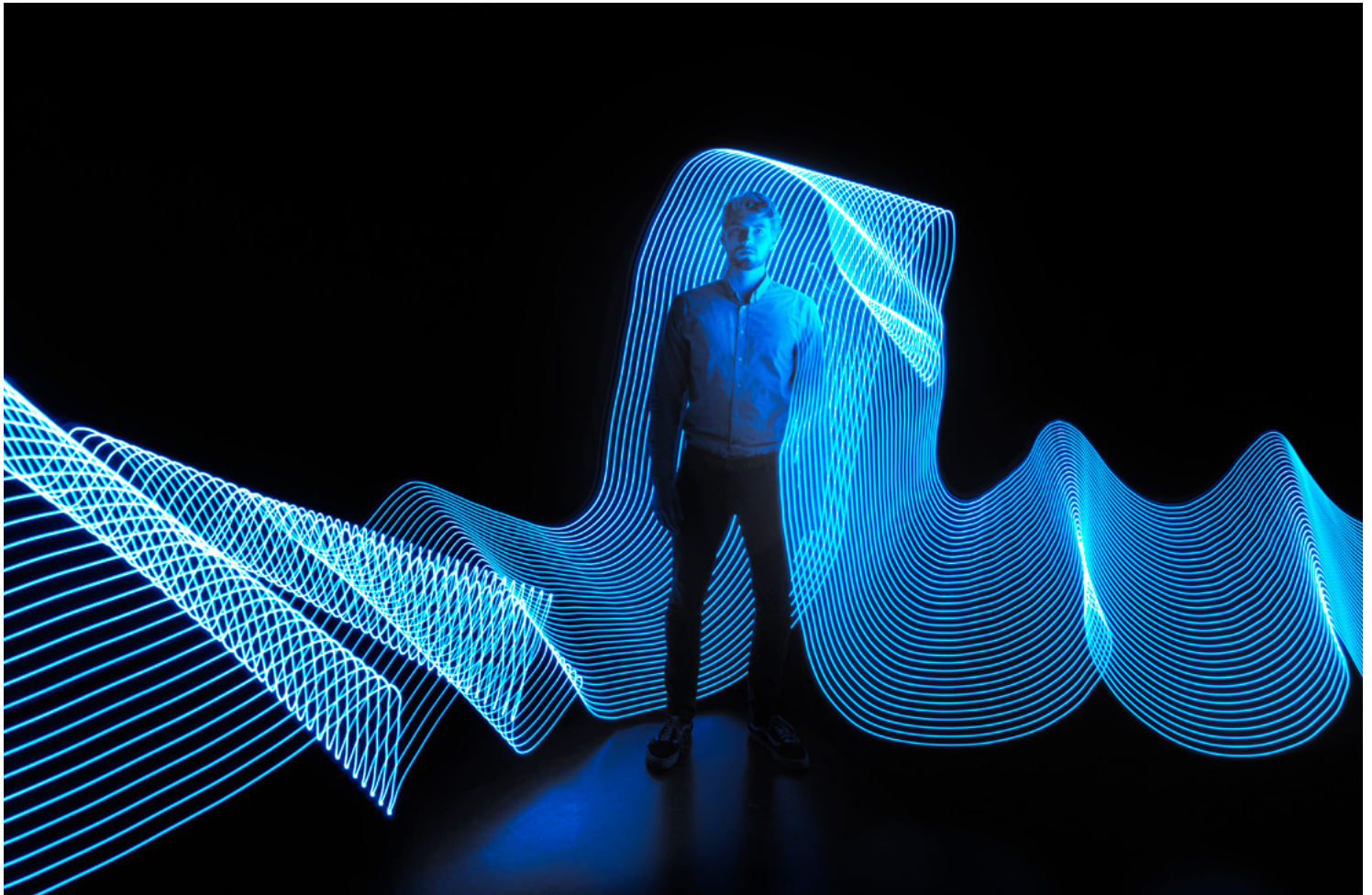
- Fatality or life-altering injury
- Catastrophic property loss
- Credible third-party claim exposure
- Credible criminal exposure
- Credible allegation of intentional misconduct
- Litigation that is reasonably anticipated

Appendix 2. Recommended language for AI outputs and reports

Use consistent labels so readers can distinguish drafts from final conclusions. Suggested labels for drafts include:

- "AI-generated working draft"
- "Candidate contributing factor"
- "Candidate corrective action"
- "Not approved as a company finding until signed off"

For final reports: 'Findings and corrective actions in this report reflect the organization's approved conclusions following human review.'



Appendix 3. References

1. Federal Rule of Civil Procedure 26(b)(1), Discovery Scope and Limits (Cornell Legal Information Institute), https://www.law.cornell.edu/rules/frcp/rule_26
2. Federal Rule of Civil Procedure 34, Producing Documents, Electronically Stored Information, and Tangible Things (Cornell Legal Information Institute), https://www.law.cornell.edu/rules/frcp/rule_34
3. Federal Rule of Civil Procedure 37(e), Failure to Make Disclosures or to Cooperate in Discovery; Sanctions, Failure to Preserve Electronically Stored Information (Cornell Legal Information Institute), https://www.law.cornell.edu/rules/frcp/rule_37
4. Federal Rule of Evidence 403, Excluding Relevant Evidence for Prejudice, Confusion, Waste of Time, or Other Reasons (Cornell Legal Information Institute), https://www.law.cornell.edu/rules/fre/rule_403
5. Federal Rule of Evidence 407, Subsequent Remedial Measures (Cornell Legal Information Institute), https://www.law.cornell.edu/rules/fre/rule_407
6. Federal Rule of Evidence 801(d)(2), Statements That Are Not Hearsay, An Opposing Party's Statement, including adoption concepts (Cornell Legal Information Institute), https://www.law.cornell.edu/rules/fre/rule_801
7. *Upjohn Co. v. United States*, 449 U.S. 383 (S. Ct. 1981), <https://supreme.justia.com/cases/federal/us/449/383/>
8. *In re Kellogg Brown & Root, Inc.*, 756 F.3d 754 (D.C. Cir. 2014), <https://law.shu.edu/documents/in-re-kellogg-brown-and-root-inc.pdf>
9. *Wal-Mart Stores, Inc. v. Wittke*, Florida Second District Court of Appeal (2016), (discussion of internal policies and the standard of care), <https://law.justia.com/cases/florida/second-district-court-of-appeal/2016/2d15-4099.html>
10. *Harris v. One Hope United, Inc.*, 2015 IL 117200 (Ill. Mar. 19, 2015) (declining to recognize a self-critical analysis privilege), <https://cases.justia.com/illinois/supreme-court/2015-117200.pdf>
11. *Hickman v. Taylor*, 329 U.S. 495 (1947) (work product doctrine origin), <https://supreme.justia.com/cases/federal/us/329/495/>
12. Restatement (Second) of Torts § 324A (negligent undertaking principles) (American Law Institute)
13. OSHA, Standards, Letters of Interpretation, Elements necessary for a violation of the General Duty Clause (letter from R. Fairfax, Directorate of Enforcement Programs to M. Racic, Int'l Bhd. Boilermakers, Dec. 18, 2003), <https://www.osha.gov/laws-regs/standardinterpretations/2003-12-18-1>
14. OSHA, Standards, Federal Employer Rights and Responsibilities Following an OSHA Inspection-1996, (definition of willful violations), <https://www.osha.gov/publications/fedrites>
15. OSHA Field Operations Manual, Chapter 4-V, Willful Violations (willful violation guidance), <https://www.osha.gov/fom/chapter-4>
16. OSHA, Standards, 29 C.F.R. 1910.119 Process safety management of highly hazardous chemicals (incident investigation requirements and resolution of recommendations), <https://www.osha.gov/laws-regs/regulations/standardnumber/1910/1910.119>
17. EPA, 40 C.F.R. 68.81 Incident Investigation (Risk Management Program, (requirement to address and resolve findings and recommendations), <https://www.ecfr.gov/current/title-40/chapter-I/subchapter-C/part-68/subpart-D/section-68.81>



