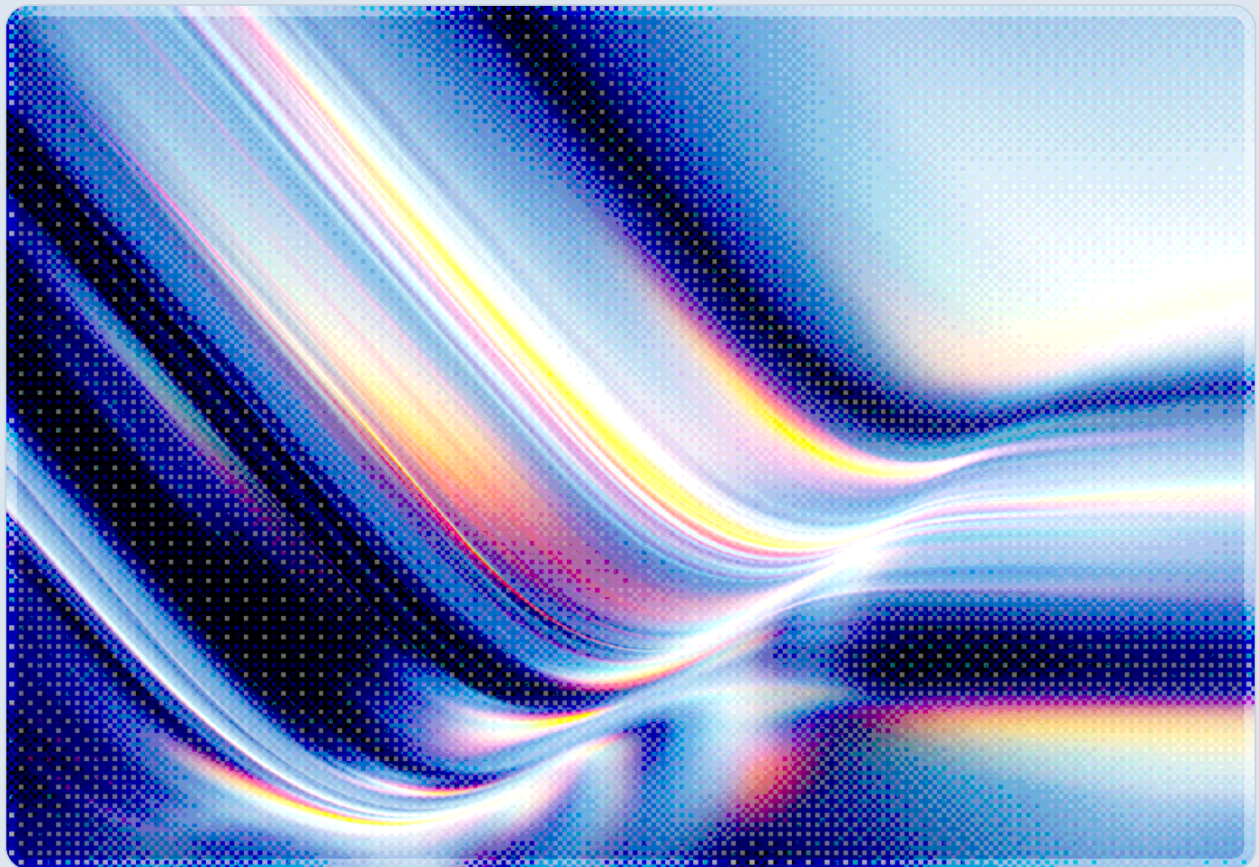


The End of Silent AI Coverage: Exclusions Outrun Risk Transfer

2026-08-30

 AI-assisted Rapid Research



CSAI

cloud
CSA security
alliance®

© 2026 Cloud Security Alliance. Some rights reserved.

You may download, store, display, view, print, redistribute, and link to this document in its original, unmodified form, provided that attribution to the Cloud Security Alliance is maintained and all trademark and copyright notices remain intact.

This document may not be modified or altered. You may quote portions of the document as permitted by the Fair Use provisions of the United States Copyright Act, provided that attribution is given to the Cloud Security Alliance.

This document may be shared on professional and social media platforms in its original form with attribution.

This document was generated with AI assistance and has not undergone official CSA review and approval processes.

Key Takeaways

- "Silent AI" – coverage that neither confirmed nor denied AI-related losses, leaving the question to be litigated at claim time – appears to be ending. Since January 1, 2026, the Insurance Services Office (ISO) has circulated three generative-AI exclusion endorsements for commercial general liability policies (CG 40 47, CG 40 48, and CG 35 08), and carrier filings adopting them have been submitted to state regulators at an accelerating pace through 2026 [1][2].
- Carrier adoption is broadening beyond CGL into the lines enterprises actually rely on for AI-driven losses: Berkley has introduced an absolute AI exclusion across directors and officers (D&O), errors and omissions (E&O), and fiduciary liability products, while Beazley and QBE have introduced AI sublimits near 10% of overall cyber policy limits rather than excluding AI outright [1][3].
- The gap this note addresses is not the existence of exclusions – CSA has documented that trend since June 2026 – but a widening lag between the exclusion wave and enterprise risk transfer assumptions: AI-related litigation volume has grown far faster than most organizations' coverage reviews, with one legal-industry tracking measure showing a 978% increase in AI-related lawsuits from 2021 to 2025 and 137% year-over-year growth from 2024 to 2025 [1].
- Four legal and disclosure events are actively shaping how the word "AI" gets read into existing policy language: Bartz v. Anthropic (a \$1.5 billion copyright settlement over training data provenance, finally approved July 20, 2026), Moffatt v. Air Canada (a 2024 tribunal ruling, now being applied at scale by 2026 courts and carriers, holding that chatbot statements bind the company issuing them), Mobley v. Workday (AI screening tools do not sever a vendor from liability for the employment decisions they influence), and Hasbro's Q1 2026 disclosure of a cyber incident – not itself AI-attributed, but illustrative of how uncertain claim recovery remains even for a well-insured enterprise – with an estimated \$60–80 million combined remediation and revenue impact [4][5].
- Underwriters willing to write affirmative AI coverage are converging on a common evidentiary bar – a documented human override capability, a human-in-the-loop inventory of consequential AI decisions, data-provenance documentation, a named accountable executive, and deepfake-resistant authentication – that overlaps substantially with existing

CSA governance frameworks, in CSA's assessment, and gives enterprises a concrete way to close the gap between what they assume is covered and what a carrier will actually pay [4].

Background

For most of the current AI deployment wave, enterprises carried AI-related exposure inside policies that were never written with AI in mind. Cyber, technology errors and omissions (Tech E&O), and commercial general liability (CGL) forms addressed AI risk only implicitly: because the policy language did not name AI as either covered or excluded, organizations and their brokers generally assumed that existing towers of coverage would respond to an AI-driven loss the same way they would respond to any other software failure. That assumption – commonly labeled "silent AI" in the surplus lines and specialty insurance press – was, in CSA's assessment, more a gap in policy drafting than a considered underwriting decision, and that gap appears to be closing rapidly through 2026 [4][6].

The Cloud Security Alliance has tracked the mechanics of that closure since mid-2026, documenting the compression of attacker timelines that is straining loss models and the January 2026 arrival of ISO's generative-AI exclusion endorsements as a structural, not temporary, shift in how AI risk gets priced. What has changed since that earlier research is not the exclusion mechanism itself but its reach and its collision with a fast-moving body of litigation. Reporting from Insurance Journal, Insurance Edge, and independent brokerage analysis published in August 2026 converges on the same finding: exclusions are migrating from a narrow set of CGL endorsements into the D&O, E&O, and cyber lines enterprises actually count on when an AI system causes harm, while the courts are simultaneously establishing that AI-mediated conduct does not insulate the company deploying it from liability [1][4][5]. This note examines that collision and what it implies for enterprises that have not yet re-audited their coverage against 2026 policy language.

Security Analysis

From CGL endorsements to the lines that matter

The three ISO endorsements that anchored the initial 2026 exclusion wave are, on close reading, narrower than their broad adoption suggests. CG 40 47 excludes coverage for bodily injury, property damage, and personal and advertising injury arising out of generative AI under the Commercial General Liability Coverage Part; CG 40 48 applies the same exclusion only to personal and advertising injury (Coverage B); and CG 35 08 applies it to the products and completed-operations hazard [1][2]. Joe

Lam, a vice president at Verisk (ISO's parent), has described the endorsements as "very essential in the marketplace," and carriers have been submitting filings adopting them to state regulators at a growing rate through 2026 [1]. Because these forms use "arising out of" rather than "caused by," carriers have interpreted them broadly: an AI system need not be the proximate cause of a loss, only part of the causal chain, for the exclusion to attach – a reading that sweeps in far more claims than a plain reading of the endorsement names would suggest [4].

CGL was never where most enterprise AI exposure actually lived, and the more consequential 2026 development is the exclusion language's spread into lines that are. Berkley has introduced an absolute AI exclusion across its D&O, E&O, and fiduciary liability products – categories that squarely cover the boards and executives whose AI governance representations are now the subject of shareholder and regulatory scrutiny [1]. Cyber carriers have generally taken a different path, introducing AI sublimits rather than outright exclusions: Beazley and QBE have each moved to cap AI-related cyber losses near 10% of the overall policy limit, converting what might once have been a fully covered loss into one covered only up to a fraction of the tower an enterprise believes it has purchased [3]. The practical effect for a risk manager is the same regardless of mechanism – exclusion or sublimit – a policy an organization has renewed for years without incident may respond to an AI-driven claim in 2026 in a materially smaller way than it did in 2025, and in some lines not at all.

Litigation is outrunning the coverage review cycle

The urgency behind this note is less about the exclusions themselves, which CSA's earlier research on the cyber insurance reckoning and the coverage gap has already analyzed in depth, than about the pace of the legal exposure those exclusions are being written against. AI-related litigation volume has grown at a rate that most enterprise legal and risk functions are not resourced to track in real time: one legal-industry measure tracking AI-related lawsuits found a 978% increase from 2021 to 2025 and a 137% year-over-year increase from 2024 to 2025, with patent infringement (11.9% of cases), copyright infringement (11.2%), and personal injury claims – including privacy violations and misuse of personal data (10.2%) – forming the largest categories [1]. That volume matters to underwriters because it converts what was, as recently as 2024, a largely theoretical tail risk into a population of resolved and pending cases carriers can use to price exclusions and sublimits – and it matters to enterprises because their existing coverage reviews are typically calibrated to an annual renewal cycle that cannot keep pace with a litigation landscape moving on a monthly one.

Four legal and disclosure events illustrate the kind of exposure now driving underwriting decisions, and each closes off a defense enterprises may have assumed was available, even though not all of them originate from 2026 itself. *Bartz v. Anthropic*, the class action brought by authors Andrea Bartz, Charles Graeber, and Kirk Wallace Johnson over the use of pirated books in training data, received final court

approval of its \$1.5 billion settlement on July 20, 2026 – the largest copyright class-action recovery on record and a direct data point for how expensive a training-data provenance failure can become [5][7] [8]. *Moffatt v. Air Canada*, decided by a Canadian tribunal in February 2024, rejected the airline's argument that its chatbot was a separate legal entity responsible for its own erroneous statements, holding that "it makes no difference whether information comes from a static page or a chatbot"; that two-year-old precedent is now the one 2026 courts and carriers are applying at scale, establishing that a company cannot disclaim liability for AI output it deploys in a customer-facing role [4]. *Mobley v. Workday* extended the same logic to employment decisions, with the court rejecting a vendor-liability defense and treating an AI hiring tool as inseparable from the human decision-making it informs, a holding directly relevant to any enterprise using AI screening or evaluation tools under the assumption that the vendor, not the deploying company, bears the resulting discrimination liability [4]. And Hasbro's first-quarter 2026 disclosure of a cyber incident – not itself AI-attributed, but illustrative of how uncertain recovery remains even for an enterprise that believed itself covered – carried an estimated \$20 million in remediation costs and \$40–60 million in delayed revenue [4]. None of these required a novel legal theory; each applied existing liability doctrine to a fact pattern that is now colliding with AI-era coverage language, which is precisely why carriers are treating the resulting exposure as a pricing problem for existing lines rather than a niche product opportunity.

What underwriters are asking for before they will write affirmative cover

A parallel and more constructive trend is the emergence of a common evidentiary standard among underwriters willing to write affirmative AI coverage rather than exclude the risk outright. One detailed account of carrier renewal conversations through mid-2026 identifies six recurring requirements: a documented human override or "kill switch" capability with a recently tested escalation path; an inventory of AI systems that make or influence decisions affecting customers, employees, or regulated outcomes; a data-provenance audit documenting the sources and labeling of training and input data; a single named, accountable executive whose signature – not a committee's – attests to AI governance controls; deepfake-resistant, out-of-band authentication for high-value transactions; and demonstrable enforcement of these controls on the technology an organization already operates, evidenced through existing logging and identity infrastructure [4]. This checklist is drawn from a single trade source and has not yet been independently corroborated across multiple carriers, but its individual elements track closely with control categories already familiar from enterprise AI governance work. Munich Re's Mosaic program, Lloyd's-backed Armilla, and newer entrants such as Counterpart and Testudo have reportedly begun writing affirmative AI coverage against evidence of these controls, while the Artificial Intelligence Underwriting Company's AIUC-1 standard – developed with input from Stanford, MIT, MITRE, and the Cloud Security Alliance [9] – ties certification directly to underwriting, insuring the customers of a certified AI agent against the specific failure modes the standard is designed to prevent [4][10].

Table 1 summarizes how the pattern differs across the lines most exposed to AI risk in 2026.

Policy Line	2026 Treatment	Representative Carriers	Enterprise Implication
Commercial General Liability	Exclusion (ISO CG 40 47/48, CG 35 08)	Broad market adoption; filings increasingly submitted for state approval	AI-linked bodily injury, property damage, and advertising injury claims presumptively unpaid
D&O / E&O / Fiduciary	Absolute exclusion	Berkley	Board and executive AI governance representations now uninsured against
Cyber	Sublimit (~10% of policy limit) rather than outright exclusion	Beazley, QBE	AI-driven breach may be covered only up to a fraction of the assumed tower
Affirmative AI products	New standalone/parametric coverage tied to control evidence	Munich Re Mosaic, Armilla (Lloyd's), Counterpart, Testudo	Coverage available, but only against documented governance evidence (kill switch, provenance audit, named accountable executive)

Why the gap persists

The underlying reason enterprise risk transfer assumptions lag the exclusion wave is structural rather than a matter of insufficient attention. Standard E&O, cyber, and general liability policies were drafted around a model of human error or discrete system breach; generative AI and agentic systems produce losses through unpredictable, probabilistic output rather than a single identifiable failure, which is difficult to fit inside policy language built for deterministic causation [6]. Brokers and risk managers conducting a conventional annual policy review are, by design, looking backward at renewal language rather than forward at a litigation landscape doubling year over year, so an exclusion or sublimit introduced mid-cycle can go unnoticed until a claim is denied. The result is a gap that will not close

through more frequent renewal reviews alone; it requires enterprises to treat AI governance evidence – the same kind of documentation underwriters are now demanding for affirmative cover – as a standing operational discipline rather than a renewal-season exercise.

Recommendations

Immediate Actions

Enterprises should request written confirmation from every carrier across their D&O, E&O, cyber, and general liability towers on whether AI-instrumented losses are covered, excluded, or sublimited under current policy language, rather than relying on broker summaries or prior-year assumptions, since the "arising out of" language in the 2026 ISO endorsements has been read broadly enough to reach claims an enterprise might not initially categorize as AI-related [1][4]. Legal and risk teams should specifically review any customer-facing AI system – chatbots, screening tools, automated decisioning – against the fact patterns in *Moffatt v. Air Canada* and *Mobley v. Workday*, since both cases establish that vendor or tooling arguments will not shield the deploying enterprise from liability [4].

Short-Term Mitigations

Organizations should build the documentation set underwriters are already asking for, regardless of whether they are actively negotiating affirmative AI coverage: a tested human override capability for consequential AI systems, an inventory of AI-influenced decisions affecting customers or employees, a data-provenance record for training and input data, and a single named executive accountable for AI governance attestations [4]. This documentation set is not incremental overhead layered on top of existing security programs; it overlaps substantially with control evidence CSA's AI Controls Matrix (AICM) v1.1 already asks organizations to maintain, and building it once serves both governance and underwriting purposes.

Strategic Considerations

Enterprises should treat the pace of AI-related litigation, not the annual renewal calendar, as the trigger for coverage review, given that legal exposure has grown at a rate – a 978% increase in AI lawsuits since 2021 – that a once-a-year policy check cannot track [1]. Boards and risk committees should also recognize that the same governance evidence increasingly required for affirmative AI insurance coverage

is now directly relevant to D&O exposure: absolute AI exclusions in D&O and fiduciary lines mean that undocumented AI governance is no longer merely a security or compliance gap, it is now also an uninsured one for the individuals making the relevant decisions [1].

CSA Resource Alignment

This note builds on the Cloud Security Alliance's ongoing analysis of AI risk insurability and the cyber insurance coverage gap, which has tracked the compression of attacker timelines straining insurer loss models and the structural, rather than temporary, nature of the January 2026 ISO exclusion endorsements. That analysis has argued that AI risk violates the conditions insurers need to price a loss cleanly – that a loss be estimable, statistically independent, and bounded – and the litigation volume documented in this note is direct evidence of the "estimable" condition failing in real time, as claim frequency and severity data become obsolete faster than actuarial models can be recalibrated. It has also promoted the practice of mapping coverage line-by-line against AI-driven loss scenarios, an approach this note's Table 1 extends with the specific carrier behavior – Berkley's absolute D&O/E&O exclusion, Beazley and QBE's cyber sublimits – observed since that earlier work.

Across this note and CSA's related research, the AI Controls Matrix (AICM) v1.1 remains the standing framework enterprises can use to build the governance evidence – human oversight, data provenance, accountable ownership – that overlaps substantially with what underwriters are asking for before writing affirmative AI coverage. No carrier has yet publicly confirmed that AICM evidence alone satisfies its underwriting bar, but the overlap gives enterprises a concrete, already-available starting point for closing the gap this note describes.

References

- [1] Insurance Journal. "[Insurer Interest in AI Coverage Exclusions Growing as Risk Becomes Omnipresent](#)." Insurance Journal, August 17, 2026.
- [2] Fenwick. "[The End of 'Silent AI'? Emerging AI Exclusions, Coverage Fragmentation, and Practical Implications for Policyholders](#)." Fenwick, 2026.
- [3] Business Insurance. "[Insurers, Brokers Adjust as AI Exclusions Emerge](#)." Business Insurance, 2026.
- [4] EPC Group. "[Silent AI Is Dead: What Six Carriers Told Me About Your 2026 Renewal](#)." EPC Group, 2026.
- [5] Authors Guild. "[Court Grants Final Approval of \\$1.5 Billion Anthropic Copyright Settlement](#)." Authors Guild, July 2026.
- [6] Insurance Edge. "[The AI Insurance Illusion: Closing the Coverage Gap Before Litigation Hits](#)." Insurance Edge, August 13, 2026.
- [7] Copyright Alliance. "[What to Know About the \\$1.5 Billion Bartz v. Anthropic Settlement](#)." Copyright Alliance, 2026.
- [8] NPR. "[Authors Have Mixed Feelings About the \\$1.5B Anthropic Copyright Infringement Ruling](#)." NPR, July 27, 2026.
- [9] Cloud Security Alliance. "[CSA Extends Leadership into Agentic AI with Addition of AIUC-1 Certification to STAR Registry](#)." Cloud Security Alliance, June 30, 2026.
- [10] Workstreet. "[What Is AIUC-1? The First Security Standard Built for AI Agents](#)." Workstreet, 2026.