

The AI Security Insurance Underwriting Framework

Standardizing Risk Assessment for the AI Era

The Underwriting Gap No One Can Afford

The global cyber insurance market has expanded to \$16 billion in 2025,² with projections reaching \$30–50 billion by 2030.⁶ Yet this rapidly growing market faces an existential gap: the AI risks that no existing framework can adequately quantify, price, or underwrite.

In November 2025, the Financial Times reported that major insurers — including Great American, Chubb, and W.R. Berkley — had formally requested regulatory permission to exclude AI liabilities from their policies.¹ By January 2026, ISO had filed absolute AI exclusions for commercial general liability and completed products/operations policies.¹ The message was unmistakable: the insurance industry does not know how to price AI risk.

This creates a paradox at the heart of the digital economy. Enterprises are told they must adopt AI or fall behind. Simultaneously, the insurance market is retreating from covering the very risks that AI adoption creates. As WTW has observed, "AI transgresses all lines of insurance — it amplifies traditional cyber risk while simultaneously introducing novel regulatory and liability exposures that traditional policy language simply was not designed to absorb."⁴

\$16B

Cyber Insurance Market 2025

83%

Lack Automated AI Controls

200+

Active AI Legal Cases

Our research reveals the depth of unpreparedness: 83% of organizations lack automated AI controls,¹² 86% have no visibility into AI data flows,¹² and 74% report only moderate or limited AI risk governance coverage.¹¹ Meanwhile, over 200 active legal cases involving AI issues — from data bias to intellectual property infringement to discrimination — are making their way through courts.¹

Current risk frameworks were not built for this challenge. NIST AI RMF provides voluntary guidance without quantitative scoring.⁸ The FAIR model depends on data quality that does not yet exist for AI systems.⁹ CIS Controls were designed for traditional cyber threats, not for LLM hallucinations or model poisoning attacks.¹⁰ None of these frameworks were built to help insurers price risk — they help organizations manage risk, which is necessary but insufficient for underwriting.

The AI Security Insurance Underwriting Framework fills this void. Developed by AI Security Intelligence (ASI), it provides a standardized, quantitative methodology for assessing AI risk across five domains and twenty-five discrete factors, producing a composite AI Insurance

Readiness Score (AIRS) that maps directly to underwriting decisions, pricing tiers, and coverage eligibility.

This white paper presents the framework in full — from the crisis driving its necessity, through the risk taxonomy it addresses, to the scoring methodology and implementation guidance for both insurers and enterprises. Our goal is not merely to describe the problem, but to offer the industry a practical, field-tested solution.

-
1. GoWalnut: <https://www.gowalnut.com/insight/the-ai-coverage-paradox-why-cyber-insurance-is-at-a-crossroads>
 2. S&P; Global: <https://www.spglobal.com/ratings/en/regulatory/article/cyber-insurance-market-outlook-2026-resilient-earnings-tougher-competition-pockets-of-growth-s101658506>
 4. WTW: <https://www.wtwco.com/en-us/insights/2025/03/emerging-ai-exposures-and-the-role-of-cyber-and-e-and-o-insurance>
 6. Risk & Insurance: <https://riskandinsurance.com/cyber-insurance-market-set-for-explosive-growth-amid-emerging-threats-and-regulatory-pressure/>
 8. NIST: <https://www.nist.gov/itl/ai-risk-management-framework>
 9. SecurityScorecard: <https://securityscorecard.com/ebook/why-the-fair-model-can-be-so-unfair/>
 10. SaferAI: <https://www.safer-ai.org/beyond-benchmarks-in-ai-cyber-risk-assessment>
 11. IBM: <https://www.ibm.com/think/insights/cios-ai-risk-governance-gap>
 12. Kiteworks: <https://www.kiteworks.com/cybersecurity-risk-management/ai-security-gap-2025-organizations-flying-blind/>

The Crisis

Chapter 1

The AI Coverage Paradox

The cyber insurance market is one of the fastest-growing segments in the global insurance industry. Global cyber premiums reached approximately \$15.3 billion at year-end 2024,² expanding to \$16 billion in 2025.² Industry projections indicate growth to \$30–50 billion by 2030,⁶ fueled by escalating digital transformation, regulatory requirements, and the ever-expanding threat landscape.

Yet beneath these growth figures lies a structural fracture. In November 2025, the Financial Times reported that several of the world's largest property and casualty insurers — including Great American, Chubb, and W.R. Berkley — had formally asked U.S. regulators for permission to exclude AI-related liabilities from their policies.¹ This was not a subtle market adjustment; it was an industry-wide acknowledgment that AI risk, as currently understood, is uninsurable under existing frameworks.

By January 2026, ISO had filed absolute AI exclusions for commercial general liability (CGL) and completed products/operations policies.¹ These exclusions are comprehensive, covering AI-generated errors and misrepresentations, hallucinated outputs, flawed chatbot advice, automated decision-making failures, and model-produced content that infringes, defames, or discriminates.¹

The Coverage Migration Problem

Exclusions do not eliminate risk — they migrate it. As Resilience Chief Underwriting Officer Maria Long has observed, exclusions push "AI exposures onto cyber and Tech E&O policies that were not designed to carry them."⁵ This creates a dangerous dynamic: enterprises believe they have coverage because AI risk has migrated to other policies, while those policies lack the specific language, limits, and underwriting rigor to properly absorb AI exposures.

The Lockton Re and Armilla AI "Ready or Not" report (February 2026) found that CGL insurers do not currently model, underwrite, or price AI risks — creating a growing gap between what insurers intend to cover and what they actually cover.¹³ Their central recommendation: AI needs its own risk class entirely, with each AI model underwritten on its individual merits — industry, foundation model, version, and use case.¹⁴

Why Insurers Cannot Price AI Risk

Five structural factors make AI risk fundamentally different from any risk class insurers have previously modeled:

1. Sparse historical loss data. AI-specific claims are too new and too varied to support actuarial modeling with statistical confidence.¹³
2. No stable legal framework. Liability doctrines are evolving rapidly, with over 200 active legal cases and new state laws creating private rights of action.¹¹⁹
3. Opaque model behavior. The "black box" nature of foundation models means insurers cannot inspect the risk-generating mechanism the way they can for physical assets or even traditional software.⁷
4. Systemic exposure from shared infrastructure. Shared foundation models (GPT-4, Claude, Gemini) create portfolio-level concentration risk. A single compromised model could trigger thousands of simultaneous losses across an insurer's entire book.¹³
5. Non-deterministic outputs. AI systems produce different outputs for the same inputs, making traditional negligence constructs — which assume predictable cause-and-effect — difficult to apply.¹⁵

Forrester projects that written cyber premiums will rise 15% in 2026 as the expansion of AI threat surfaces reverses the recent market softening.⁵ The question is whether that premium growth will be accompanied by frameworks capable of actually quantifying the risks being priced.

1. GoWalnut: <https://www.gowalnut.com/insight/the-ai-coverage-paradox-why-cyber-insurance-is-at-a-crossroads>

2. S&P; Global: <https://www.spglobal.com/ratings/en/regulatory/article/cyber-insurance-market-outlook-2026-resilient-earnings-tougher-competition-pockets-of-growth-s101658506>

5. Carrier Management: <https://www.carriermanagement.com/features/2026/03/09/285417.htm>

6. Risk & Insurance:

<https://riskandinsurance.com/cyber-insurance-market-set-for-explosive-growth-amid-emerging-threats-and-regulatory-pressure/>

7. Swiss Re: <https://www.swissre.com/institute/research/sonar/sonar2024/ai-silent-cyber.html>

13. Lockton Re / Armilla AI:

<https://global.lockton.com/re/en/news-insights/ready-or-not-the-impact-of-artificial-intelligence-on-insurance-risks>

14. Carrier Mgmt – AI Risk Class: <https://www.carriermanagement.com/news/2026/02/17/284646.htm>

15. IAPP: <https://iapp.org/news/a/how-ai-liability-risks-are-challenging-the-insurance-landscape>

19. Wiley: <https://www.wiley.law/article-12233>

Why Current Frameworks Fall Short

Three established frameworks dominate risk assessment and cybersecurity governance today. Each offers valuable contributions, but none was designed to help insurers price AI-specific risk.

NIST AI Risk Management Framework (AI RMF 1.0)

The NIST AI RMF, along with its GenAI Profile (AI 600-1, July 2024), provides a structured approach to AI risk governance through four functions: Govern, Map, Measure, and Manage.⁸ It is the most comprehensive government-sponsored framework available. However, it is voluntary, has seen incomplete market adoption, and provides no insurance-specific guidance, no quantitative risk scoring, and no actuarial methodology. Implementation is resource-intensive, making it particularly challenging for mid-size organizations that represent a significant portion of the insured market.

FAIR Model (Factor Analysis of Information Risk)

FAIR offers a quantitative approach to information risk analysis and has gained traction in the cybersecurity community. However, as SecurityScorecard has documented, the model is heavily dependent on data quality and timeliness — two attributes that AI risk data fundamentally lacks.⁹ The model contains no specific variables for loss magnitude controls, and many practitioners resort to subject-matter-expert-estimated data ranges, introducing subjectivity into what is supposed to be a quantitative process. FAIR was not designed for probabilistic AI systems and offers no native support for the non-deterministic, interconnected nature of AI risk.

CIS Controls

CIS Controls represent well-established security best practices, but they were designed for traditional cybersecurity threats. An AWS survey found that approximately 40% of leaders now rank AI-based frameworks as the top lever for reducing cyber risk,¹⁰ signaling that the market recognizes the gap between traditional controls and AI-specific threats. Critical gaps exist in LLM abuse detection, model input/output monitoring, and AI data flow governance.

Key Insight: None of these frameworks were built to help insurers price risk. They help organizations manage risk — a necessary but insufficient condition for underwriting. The insurance industry needs a framework that translates risk management maturity into actuarial inputs: quantitative scores, weighted domains, and tier classifications that map directly to coverage decisions and premium calculations.

8. NIST: <https://www.nist.gov/itl/ai-risk-management-framework>

9. SecurityScorecard: <https://securityscorecard.com/ebook/why-the-fair-model-can-be-so-unfair/>

The AI Risk Taxonomy

Chapter 3

The Five Domains of AI Insurance Risk

Our research team has developed a proprietary risk taxonomy that organizes the full spectrum of AI insurance exposures into five distinct domains. Each domain captures a fundamentally different category of risk, with its own loss drivers, detection challenges, and actuarial implications. Together, they provide a comprehensive map of the AI insurance risk landscape.

DOMAIN 1

Model Integrity Risk

CRITICAL

Data poisoning, adversarial attacks, model manipulation, and training data compromise.

200–400 poisoned samples can compromise models | Federated learning amplifies risk

Model integrity risk encompasses threats that compromise the fundamental reliability of AI systems at the model level. Research published in JMIR demonstrates that as few as 200–400 poisoned training samples can compromise models trained on millions of data points, with attack success depending on absolute sample count rather than poisoning rate.¹⁷ In healthcare contexts, compromising AI diagnostics can cost as little as \$20,000–\$200,000.¹⁷ Federated learning architectures, increasingly popular for privacy-preserving AI, actually amplify poisoning risks while simultaneously hindering detection. Byzantine-robust aggregation methods remain inadequate against sophisticated attack strategies.²⁹

The adversarial attack surface is expanding rapidly. IBM X-Force reports that over 80% of phishing attempts are now AI-generated, and these attacks are four times more likely to deceive recipients than their human-crafted counterparts.³ The FBI IC3 logged 193,000 phishing and spoofing complaints in 2024, with wire fraud exceeding \$109 million.³

DOMAIN 2

Output Liability Risk

HIGH

Hallucinations, discriminatory outputs, misguidance, and downstream liability from AI-generated content.

200+ active legal cases | API providers disclaim responsibility

Output liability risk arises when AI systems generate content that causes harm — through hallucinated facts, discriminatory recommendations, or misleading guidance. The legal landscape is evolving rapidly, with over 200 active cases involving AI-related claims spanning data bias, intellectual property infringement, and discrimination.¹ The Lockton Re report includes a particularly instructive claim scenario: an AI chatbot generates incorrect warranty commitments without any breach or malicious activity — a pure output liability event that traditional cyber policies would not cover.¹³

A critical structural issue compounds this risk: major API providers — including OpenAI, Anthropic, and Google — disclaim responsibility for the outputs their models generate, effectively pushing liability downstream to deployers and end users.¹⁶ Healthcare, legal, and financial services face the highest output liability exposure.

1. GoWalnut: <https://www.gowalnut.com/insight/the-ai-coverage-paradox-why-cyber-insurance-is-at-a-crossroads>

3. ISACA: <https://www.isaca.org/resources/news-and-trends/isaca-now-blog/2025/cyber-insurance-in-crisis-with-ai-blind-spots>

13. Lockton Re / Armilla AI:

<https://global.lockton.com/re/en/news-insights/ready-or-not-the-impact-of-artificial-intelligence-on-insurance-risks>

16. TechAndMediaLaw: <https://techandmedialaw.com/ai-hallucination-liability/>

17. PMC/JMIR: <https://pmc.ncbi.nlm.nih.gov/articles/PMC12881903/>

29. Coverlink:

<https://coverlink.com/cyber-liability-insurance/cyber-solutions-defending-ai-systems-from-malicious-data-poisoning-attacks-2/>

DOMAIN 3

Supply Chain Concentration Risk

HIGH

Third-party model dependencies, shared infrastructure, training data provenance, and vendor liability gaps.

Foundation model concentration | "Fruit of the poisonous tree" liability

The AI supply chain represents aggregate risk from every entity, dataset, and model contributing to a given AI output.¹⁸ Unlike traditional software supply chains, AI supply chains are opaque by nature — training data provenance is impossible to fully verify, and regulatory spillover means that if a foundation model provider violates the EU AI Act, every deployer inherits potential liability.¹⁸

The "fruit of the poisonous tree" doctrine presents a particularly acute supply chain risk: if training data was collected illegally or without proper consent, every model trained on that data — and every downstream application built on those models — carries embedded liability.¹⁸ This creates a liability cascade that is almost impossible to quantify using current actuarial methods.

DOMAIN 4

Regulatory Exposure Risk

ELEVATED

EU AI Act compliance, U.S. state law exposure, cross-border requirements, and enforcement penalties.

EU AI Act: fines up to 7% global turnover | Multiple U.S. states creating private rights of action

The regulatory landscape for AI is expanding at an unprecedented pace. The EU AI Act's high-risk provisions become enforceable on August 2, 2026, classifying insurance AI — including claims processing, underwriting, pricing, and fraud detection — as high-risk under Annex III Category 5(b).²¹ Compliance requires explainability layers, bias testing, decision audit trails, and human oversight. Violations of banned AI practices carry fines of up to €35 million or 7% of global turnover.²⁰

In the United States, multiple states are expanding AI liability with private rights of action. California has enacted chatbot safety protocols with penalties of \$1,000 or more per violation plus attorney fees. Texas imposes AI harm penalties of \$10,000–\$200,000, plus \$2,000–\$40,000 per day for continued violations.¹⁹ This patchwork of state laws creates a compliance surface that is expensive to navigate and impossible to fully predict.

DOMAIN 5

Systemic Correlation Risk

CRITICAL

Shared model infrastructure creating portfolio-level concentration, cascading failures, and correlated losses.

Single model failure = thousands of simultaneous claims | Traditional diversification fails

Systemic correlation risk is perhaps the most consequential domain for the insurance industry because it challenges the fundamental assumption underlying portfolio diversification: that losses across policyholders are largely independent. Shared foundation models — GPT-4, Claude, Gemini, and others — create concentration risk at the portfolio level.¹³ A single compromised model could trigger thousands of simultaneous losses across an insurer's entire book.

As the Lockton Re report observes, insurers can absorb a \$400 million loss to one company; they cannot absorb \$400 million distributed across their entire portfolio simultaneously.¹³ Compromised training data can cause failures across multiple organizations regardless of geography or industry.³⁰ This risk class is fundamentally different from anything insurers have previously modeled — including natural catastrophe risk, where geographic diversification provides meaningful protection.

13. Lockton Re / Armilla AI:

<https://global.lockton.com/re/en/news-insights/ready-or-not-the-impact-of-artificial-intelligence-on-insurance-risks>

18. TrustArc: <https://trustarc.com/resource/ai-supply-chain-risk-vendor-due-diligence/>

19. Wiley: <https://www.wiley.law/article-12233>

20. Softermii: <https://www.softermii.com/blog/artificial-intelligence/eu-ai-act-compliance-guide>

21. Munich Re – EU AI Act:

<https://www.munichre.com/automation-solutions/en/resources/insights/ai/new-eu-act-regulates-ai-in-insurance.html>

30. Verisk: <https://core.verisk.com/Insights/Emerging-Issues/Articles/2025/January/Week-4/Poisoned-Data-Represents-an-AI-Risk>

The Framework

Chapter 4

The ASI AI Security Insurance Underwriting Framework

The ASI AI Security Insurance Underwriting Framework is a 5-domain, 25-factor scoring methodology designed to provide insurers with a standardized, quantitative basis for AI risk assessment. Each domain contains five assessment factors, scored on a 1–5 scale and weighted by domain criticality. The result is a composite AI Insurance Readiness Score (AIRS) from 0 to 100 that maps directly to underwriting decisions.

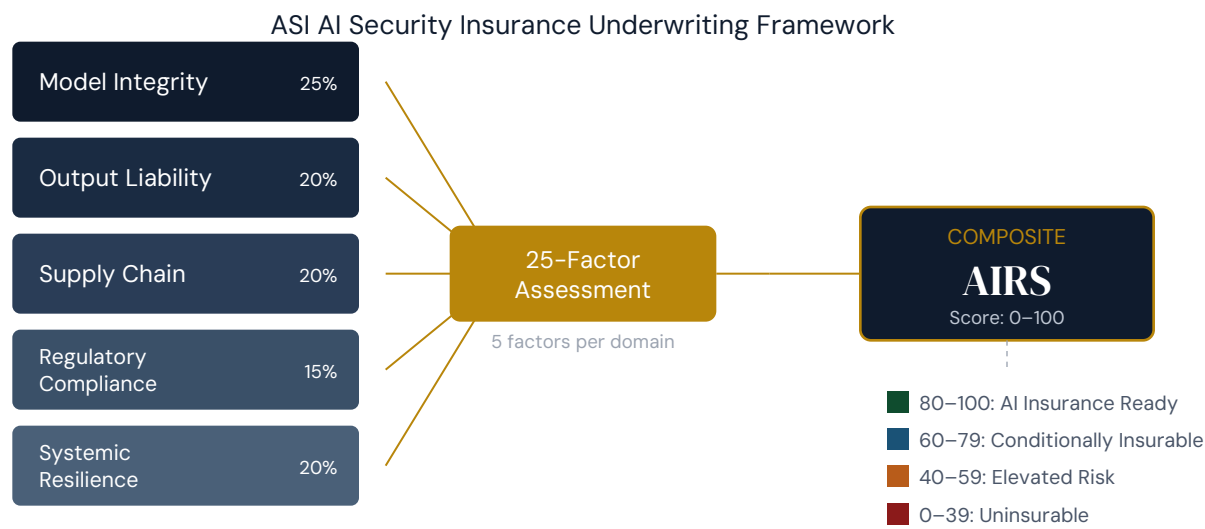


Figure 1: ASI Framework Architecture — Five Domains to Composite AIRS

Domain 1: Model Integrity — Weight: 25%

Factor	Description
1.1 Training Data Provenance	Documentation of training data sources, lineage, licensing, and consent verification for all data used in model development.
1.2 Adversarial Robustness Testing	Frequency and rigor of adversarial testing protocols, including red-teaming cadence and vulnerability remediation timelines.

1.3 Model Versioning & Rollback	Capabilities for maintaining model version history, rapid rollback to known-good states, and change management procedures.
1.4 Poisoning Detection	Mechanisms for detecting data poisoning attacks, anomalous training inputs, and compromised data pipelines.
1.5 Drift Detection & Monitoring	Continuous monitoring for model performance degradation, concept drift, and distribution shift in production environments.

Domain 2: Output Liability — Weight: 20%

Factor	Description
2.1 Human-in-the-Loop Oversight	Protocols for human review and approval of high-stakes AI outputs, escalation triggers, and override mechanisms.
2.2 Output Validation	Automated and manual mechanisms for fact-checking, accuracy verification, and hallucination detection in AI-generated content.
2.3 Disclaimer & Disclosure	Frameworks for transparent disclosure of AI involvement, output limitations, and appropriate use guidance.
2.4 Error Correction & Recall	Procedures for correcting erroneous AI outputs, notifying affected parties, and documenting remediation actions.
2.5 Liability Documentation	Contractual liability allocation in terms of service, customer agreements, and vendor contracts governing AI outputs.

Domain 3: Supply Chain — Weight: 20%

Factor	Description
3.1 Vendor Due Diligence	Thoroughness of security, compliance, and risk assessment for foundation model providers and AI vendors.
3.2 Dependency Mapping	Comprehensive inventory of third-party models, APIs, and AI components with dependency chain documentation.
3.3 Sub-processor Auditing	Auditing protocols for data enrichment partners, sub-processors, and downstream data handlers in the AI pipeline.
3.4 Contractual Liability	Contractual allocation of liability with AI vendors including indemnification, SLAs, and breach notification requirements.
3.5 Data Provenance Verification	Mechanisms for verifying training data provenance, licensing compliance, and consent chains across the supply chain.

Domain 4: Regulatory Compliance — Weight: 15%

Factor	Description
4.1 EU AI Act Readiness	Assessment of preparedness for EU AI Act high-risk requirements including conformity assessment, documentation, and human oversight.
4.2 U.S. State Law Compliance	Mapping of compliance posture across relevant U.S. state AI regulations, including California, Texas, and Colorado.
4.3 Cross-Border Data Transfer	Protocols for lawful cross-border data transfers including GDPR adequacy, standard contractual clauses, and data localization.
4.4 Bias Testing & Fairness	Documentation and cadence of algorithmic bias testing, fairness metrics, and disparate impact analysis.
4.5 Regulatory Reporting	Procedures for regulatory incident notification, mandatory reporting compliance, and records retention.

Domain 5: Systemic Resilience — Weight: 20%

Factor	Description
5.1 Model Diversification	Strategy for reducing foundation model concentration including multi-vendor approaches and fallback architectures.
5.2 Failure Isolation	Architecture for containing AI failures including circuit breakers, graceful degradation, and blast radius limitation.
5.3 Business Continuity	Business continuity and disaster recovery planning specific to AI system failures, including manual fallback procedures.
5.4 Concentration Monitoring	Ongoing monitoring of shared infrastructure dependencies, single points of failure, and concentration risk metrics.
5.5 AI Incident Response	Incident response playbooks specific to AI failures including model quarantine, output recall, and stakeholder notification.

Scoring Methodology and Risk Tiers

The AIRS composite score is calculated through a structured methodology that converts qualitative assessments into quantitative, actuarially useful outputs. Each of the 25 factors is scored on a 1–5 scale with clear, observable criteria at each level.

Scoring Rubric

Score	Rating	Definition
1	Non-Existent	No documented controls, processes, or awareness. The organization has not addressed this factor.
2	Ad Hoc	Informal or inconsistent practices exist. Controls are reactive, undocumented, and dependent on individual knowledge.
3	Defined	Documented policies and procedures exist. Controls are implemented but may not be consistently enforced or measured.
4	Managed	Controls are consistently applied, measured, and reviewed. Regular testing and improvement cycles are in place.
5	Optimized	Industry-leading practices with continuous improvement, automated monitoring, and proactive threat anticipation.

Composite Score Calculation

The AIRS composite score is calculated as follows:

$$\text{AIRS} = (\text{Model Integrity Avg} \times 0.25 + \text{Output Liability Avg} \times 0.20 + \text{Supply Chain Avg} \times 0.20 + \text{Regulatory Compliance Avg} \times 0.15 + \text{Systemic Resilience Avg} \times 0.20) \times 20$$

Each domain average is the arithmetic mean of its five factor scores (1–5). The weighted sum is multiplied by 20 to produce a 0–100 composite score. This approach ensures that a perfect score of 5 across all factors yields AIRS = 100, while a minimum score of 1 across all factors yields AIRS = 20 — reflecting that even organizations with the weakest controls have some baseline infrastructure in place.

AI Insurance Readiness Score (AIRS) – Risk Tiers

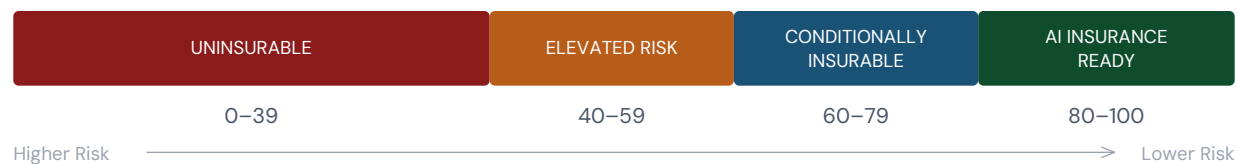


Figure 2: AIRS Risk Tier Classification

Tier	AIRS Range	Classification	Underwriting Implication
Tier 1	80–100	AI Insurance Ready	Qualifies for affirmative AI coverage with premium discounts of 10–25%. Eligible for broadest coverage terms.
Tier 2	60–79	Conditionally Insurable	Coverage available with specific exclusions, sub-limits, and higher premiums. Remediation milestones may unlock broader terms.
Tier 3	40–59	Elevated Risk	Limited coverage only, with significant exclusions and mandatory remediation requirements. Premiums reflect elevated exposure.
Tier 4	0–39	Uninsurable	Coverage denied until remediation milestones achieved. Organization provided with roadmap to reach Tier 3 minimum.

Implementation

Chapter 6

For Insurers: Integrating the Framework

The AIRS framework is designed for direct integration into existing underwriting workflows. Our analysts have developed implementation guidance based on nearly two decades of field experience in cybersecurity assessment and insurance risk evaluation.

Underwriting Questionnaire Integration

The 25-factor assessment matrix can be incorporated into existing cyber insurance application forms as a supplementary AI risk module. Each factor maps to specific, observable evidence that underwriters can request and verify: documentation artifacts, system configurations, testing reports, and governance records. This evidence-based approach reduces reliance on self-attestation and provides underwriters with auditable inputs.

Pricing Adjustments

AIRS tiers map to pricing multipliers that can be applied to base cyber premiums. Tier 1 organizations (AIRS 80–100) qualify for premium discounts reflecting their demonstrably lower risk profile. Tier 2 (AIRS 60–79) organizations receive standard pricing with exclusion-specific surcharges. Tier 3 (AIRS 40–59) organizations face elevated premiums with sub-limits on AI-specific claims. Tier 4 (AIRS 0–39) organizations are declined until remediation brings them to Tier 3 minimum thresholds.

Portfolio-Level Concentration Analysis

The Systemic Resilience domain (Domain 5) provides data that enables portfolio-level concentration analysis — a capability that current underwriting processes lack entirely.¹³ By aggregating Domain 5 scores across the book, insurers can identify portfolio-level exposure to specific foundation models, shared infrastructure providers, and correlated failure scenarios. This enables informed reinsurance purchasing and aggregate limit management.

Reinsurance Implications

Reinsurers face the same pricing challenge as primary carriers, magnified by the aggregation of AI risk across multiple cedants. The AIRS framework provides reinsurers with a standardized vocabulary and scoring methodology for evaluating cedant portfolios. Treaty terms can

For Enterprises: Improving Your AIRS

Organizations seeking to improve their insurability — or reduce their premiums — can use the AIRS framework as a self-assessment tool and remediation roadmap. Our analysts recommend a phased approach that prioritizes high-impact, achievable improvements.

Self-Assessment Methodology

Begin with an honest baseline assessment using the 25-factor matrix. Score each factor based on current state, not aspirational targets. Document evidence supporting each score. Identify the three domains with the lowest average scores — these represent the highest-impact improvement opportunities.

Quick Wins: Highest-Impact Improvements

Our analysis of organizational readiness data indicates that the fastest improvements typically come from three areas:¹²

- Output Liability (Domain 2): Implementing human-in-the-loop protocols, output disclaimers, and error correction procedures. These are process and documentation changes that can be deployed within weeks.
- Regulatory Compliance (Domain 4): Conducting EU AI Act readiness assessments and U.S. state law compliance mapping. Much of this work is documentation and gap analysis that leverages existing legal and compliance teams.
- Supply Chain (Domain 3): Creating a third-party AI model inventory and initiating vendor due diligence. Many organizations use third-party AI without a complete inventory — simply documenting dependencies raises the score.

30-60-90 Day Remediation Roadmap

Timeline	Focus Areas	Target Outcome
Days 1–30	AI asset inventory, output disclaimer frameworks, baseline AIRS self-assessment, vendor dependency mapping	Complete visibility into AI footprint; baseline AIRS established
Days 31–60	Human-in-the-loop protocols, bias testing initiation, incident response playbooks, regulatory gap analysis	Critical controls deployed; regulatory exposure mapped
Days 61–90	Adversarial robustness testing, vendor due diligence completion, model diversification planning, continuous monitoring deployment	Measurable AIRS improvement; readiness for insurer assessment

The Path Forward

Chapter 8

Building the Standard

The AIRS framework is not intended to be a proprietary tool that sits behind a single organization's paywall. Our vision at ASI is to serve as a neutral, authoritative standard-setter — fulfilling for AI insurance risk what NIST provides for cybersecurity governance and what FAIR provides for information risk quantification. The industry needs a common language, and AIRS provides it.

The Adoption Flywheel

Standardization creates a virtuous cycle. As more insurers adopt the AIRS framework, assessment data accumulates, enabling the development of actuarial models grounded in empirical loss correlations. Better actuarial models enable more accurate pricing. More accurate pricing enables broader coverage. And broader coverage generates more assessment data, further refining the framework.

- Framework Adoption — Insurers integrate AIRS into underwriting. Enterprises pursue AIRS assessments.
- Data Collection — Assessment scores correlated with loss experience across the insured portfolio.
- Actuarial Modeling — Empirical loss data enables factor-level predictive models for AI risk.
- Pricing Precision — Refined models enable granular, defensible pricing of AI exposures.
- Coverage Expansion — Confident pricing enables affirmative AI coverage products to scale.
- Framework Refinement — Market feedback and loss data inform framework updates and new factors.

A Call for Industry Collaboration

Building a durable standard requires participation from every constituency in the AI insurance ecosystem. We invite insurers and reinsurers to pilot the AIRS framework in their underwriting processes and contribute anonymized loss data to the actuarial commons. We call on regulators to recognize AIRS as an acceptable basis for demonstrating AI risk governance compliance. And we encourage enterprises to adopt AIRS self-assessment as a component of their AI governance

programs.

AI Insurance as an Enabler of Responsible Innovation

The historical parallel is instructive. Cyber insurance — once a niche product that many believed was unviable — became a critical enabler of e-commerce and digital transformation.²³ Organizations adopted better security practices not only because regulations required it, but because their insurers incentivized it through premium discounts and coverage requirements. Insurance became a market-driven mechanism for improving cybersecurity hygiene across the economy.

AI insurance can serve the same function. When organizations know that robust AI governance — measured through frameworks like AIRS — directly translates to better coverage and lower premiums, they invest in those controls. The insurance market becomes an accelerant for responsible AI adoption, not a barrier to it.²⁴

"The question is not whether AI will create systemic risk events, but when — and whether underwriting practices can keep pace."

— Lockton Re / Armilla AI, "Ready or Not" (February 2026)

The AIRS framework is our contribution to bridging the gap between AI innovation and the insurance industry's need for quantifiable, auditable risk assessment. The technology will continue to evolve. The regulatory landscape will continue to shift. But the need for a standardized, evidence-based approach to AI insurance underwriting is permanent — and the time to build it is now.

13. Lockton Re / Armilla AI:

<https://global.lockton.com/re/en/news-insights/ready-or-not-the-impact-of-artificial-intelligence-on-insurance-risks>

23. Tech Policy Press: <https://techpolicy.press/why-insurers-may-unwittingly-become-ai-safety-champions>

24. Deloitte: <https://www.deloitte.com/us/en/insights/industry/financial-services/risk-insurance-for-ai.html>

About AI Security Intelligence

AI Security Intelligence (ASI) is an independent research and advisory firm specializing in the intersection of artificial intelligence, cybersecurity, and insurance risk. Our analysts bring nearly two decades of field experience in security operations, risk assessment, and policy analysis.

Our research team holds credentials including CISSP, AI Security, AI Ethics, and Blockchain Security certifications. We provide field-tested intelligence from operationally certified analysts — not theoretical frameworks developed in isolation from the real-world threat landscape.

ASI's mission is to close the gap between AI capability and AI insurability. We work with insurers, reinsurers, enterprises, and regulators to develop the standards, methodologies, and market intelligence needed to underwrite AI risk with confidence.

Contact

Email: contact@aisecurityintelligence.com

Web: aisecurityintelligence.com

Sources

1. GoWalnut: <https://www.gowalnut.com/insight/the-ai-coverage-paradox-why-cyber-insurance-is-at-a-crossroads>
2. S&P; Global: <https://www.spglobal.com/ratings/en/regulatory/article/cyber-insurance-market-outlook-2026-resilient-earnings-tougher-competition-pockets-of-growth-s101658506>
3. ISACA: <https://www.isaca.org/resources/news-and-trends/isaca-now-blog/2025/cyber-insurance-in-crisis-with-ai-blind-spots>
4. WTW: <https://www.wtwco.com/en-us/insights/2025/03/emerging-ai-exposures-and-the-role-of-cyber-and-e-and-o-insurance>
5. Carrier Management: <https://www.carriermanagement.com/features/2026/03/09/285417.htm>
6. Risk & Insurance: <https://riskandinsurance.com/cyber-insurance-market-set-for-explosive-growth-amid-emerging-threats-and-regulatory-presures/>
7. Swiss Re: <https://www.swissre.com/institute/research/sonar/sonar2024/ai-silent-cyber.html>
8. NIST: <https://www.nist.gov/itl/ai-risk-management-framework>
9. SecurityScorecard: <https://securityscorecard.com/ebook/why-the-fair-model-can-be-so-unfair/>
10. SaferAI: <https://www.safer-ai.org/beyond-benchmarks-in-ai-cyber-risk-assessment>
11. IBM: <https://www.ibm.com/think/insights/cios-ai-risk-governance-gap>
12. Kiteworks: <https://www.kiteworks.com/cybersecurity-risk-management/ai-security-gap-2025-organizations-flying-blind/>
13. Lockton Re / Armilla AI: <https://global.lockton.com/re/en/news-insights/ready-or-not-the-impact-of-artificial-intelligence-on-insurance-risks>
14. Carrier Mgmt – AI Risk Class: <https://www.carriermanagement.com/news/2026/02/17/284646.htm>
15. IAPP: <https://iapp.org/news/a/how-ai-liability-risks-are-challenging-the-insurance-landscape>
16. TechAndMediaLaw: <https://techandmedialaw.com/ai-hallucination-liability/>
17. PMC/JMIR: <https://pmc.ncbi.nlm.nih.gov/articles/PMC12881903/>
18. TrustArc: <https://trustarc.com/resource/ai-supply-chain-risk-vendor-due-diligence/>
19. Wiley: <https://www.wiley.law/article-12233>
20. Softermii: <https://www.softermii.com/blog/artificial-intelligence/eu-ai-act-compliance-guide>
21. Munich Re – EU AI Act: <https://www.munichre.com/automation-solutions/en/resources/insights/ai/new-eu-act-regulates-ai-in-insurance.html>
22. Debevoise: <https://www.debevoise.com/insights/publications/2025/05/europes-regulatory-approach-to-ai-in-the-insurance>
23. Tech Policy Press: <https://techpolicy.press/why-insurers-may-unwittingly-become-ai-safety-champions>
24. Deloitte: <https://www.deloitte.com/us/en/insights/industry/financial-services/risk-insurance-for-ai.html>
25. Munich Re – Cyber Gap: <https://www.munichre.com/en/insights/cyber/cyber-protection-gap.html>
26. Mordor Intelligence: <https://www.mordorintelligence.com/industry-reports/ai-in-insurance-market>
27. Precedence Research: <https://www.precedenceresearch.com/artificial-intelligence-in-insurance-market>
28. Risk Strategies: <https://www.risk-strategies.com/state-of-the-insurance-market-2025-outlook-cyber>
29. Coverlink: <https://coverlink.com/cyber-liability-insurance/cyber-solutions-defending-ai-systems-from-malicious-data-poisoning-attacks-2/>
30. Verisk: <https://core.verisk.com/Insights/Emerging-Issues/Articles/2025/January/Week-4/Poisoned-Data-Represents-an-AI-Risk>