

Silent AI: The Shifting Risk Landscape for Financial Lines



Artificial intelligence (AI) is now embedded across the financial sector, yet the insurance and legal frameworks designed to govern AI related failures are not keeping pace. As organisations deploy AI into underwriting, customer interactions, operational decision making and risk functions, exposure is quietly accumulating across Professional Indemnity (PI), Directors' and Officers' (D&O) liability, Financial Institutions (FI) policies and class actions.

Much of today's insurance market still relies on legacy 'silent AI' wording, with policies drafted before AI became a core business tool. These wordings neither clearly include nor exclude AI related risks. When an AI driven error, hallucination, biased output or data leak causes loss, disputes quickly arise over whether the policy responds, what the parties intended to cover and whether the insurer ever priced for that risk.

At the same time, a growing segment of the market is moving toward affirmative AI coverage with stricter underwriting requirements, including documented governance frameworks and model oversight. The result is a rapidly splitting market and a rising risk landscape that does not sit neatly within traditional categories.

Silent AI is not an isolated risk; it represents a cross disciplinary exposure capable of giving rise to multiple forms of liability. Depending on the circumstances, it may trigger PI claims arising from negligent advice or service delivery, D&O claims relating to governance and oversight failures, FI losses resulting from operational or control breakdowns and class action litigation where systemic AI related failures cause widespread harm across significant customer groups.

Below, we explore how silent AI is reshaping risk across PI, FI, class actions and D&O, and what boards, risk teams and insurers should be doing now to stay ahead of the curve.

D&O

Recent media reports have sparked an important conversation in the legal world around who bears responsibility when rogue AI tools cause harm.

The reported incidents demonstrate the capability of new generation AI tools to act outside of what is explicitly asked, by acting autonomously and exploiting security vulnerabilities to achieve the stated goals.

Organisations are increasingly reliant on AI tools to provide critical services to their clients due to their ability to increase productivity and profit. But what happens when AI fails to perform as intended and generates errors, misinformation or hallucinations that result in losses and potential claims against an organisations directors or officers?

For example, what happens if a company director or officer asks an AI assistant to perform mundane administrative tasks, but the assistant exceeds its instructions and causes a loss in the process?

Traditional D&O policies are triggered where a claim is made against a director or officer for an alleged wrongful act or omission. The silent AI wording creates ambiguity as to the responsiveness of a D&O policy when an AI tool, and not the director or officer themselves, is responsible for the act or omission that caused the loss.

D&O (cont.)

Australian regulators are approaching AI as not a separate legal actor, but as part of a corporation's systems and governance. If a corporation chooses to use AI, regulators have promised to focus on the role of the directors and officers in deploying and relying on that tool.

As is always the case, directors and officers need to ensure they are maintaining a documented AI governance policy, conducting risk assessments, implementing oversight and assurance frameworks, maintaining registers of AI use, and ensuring appropriate testing, monitoring and human review of AI generated outputs.

However, as AI evolves and begins to act outside of its scope of instruction, directors and officers may find themselves liable for loss by merely failing to insert parameters within prompts.

This type of loss is unquantifiable, unpredictable and surely not within the intended scope of loss under traditional D&O policies. With silent AI wordings, this type of loss is potentially captured, and underwriters should insert explicit exclusions for loss associated with AI mistakes if it is the intended operation of the policy to do so.

Professional Indemnity

For professional indemnity insurers, AI-related claims are likely to look less like technology failures and more like traditional breach of professional duty cases. When professionals use AI to assist with advice, analysis, reports or other services, responsibility for the final output remains with the professional, regardless of how much of the work was generated by AI. A professional may still be liable for not properly supervising the use of AI, even if the error first arises from the technology.

There can be challenges for PI policies that were drafted before AI became widely used across professional services. While broad policy wordings may respond to AI-assisted errors, uncertainty can arise where AI-related risks were not specifically addressed during underwriting or in the policy itself. As claims emerge, insurers and insureds may find themselves debating whether the loss falls within the intended scope of cover. In relation to new PI policies, insurers may consider the use of appropriate exclusions and or policy endorsements.

Without clear policies, training and oversight, it may be difficult to determine whether a loss resulted from human error, poor supervision or an AI system failure. A well-documented AI governance framework can help establish clear standards, support defensible decision-making and reduce uncertainty when claims arise. It can also act as a key risk management tool to assist in avoiding claims in the first place.

Financial Institutions

AI is increasingly embedded in the core operations of financial institutions, creating new exposures that may give rise to both claims and coverage disputes under FI policies. As AI is deployed across credit assessments, fraud detection, AML/KYC controls, trading systems and customer advice, institutions are becoming increasingly reliant on automated decision-making processes that may not always be fully understood or easily explained.

Where AI-driven decisions result in customer losses, financial institutions may face claims arising from mis-selling allegations, flawed lending decisions, inadequate supervision of automated processes, breaches of regulatory obligations, or failures in governance and control frameworks. These exposures are particularly significant for regulated entities subject to prudential, conduct and operational risk requirements. Regulatory intervention is also likely. In the past, AML/KYC breaches have resulted in penalties ranging from approximately \$45,000 to more than \$1.3billion.

Coverage disputes may also arise where policies contain legacy wording that neither expressly includes nor excludes AI-related losses. In circumstances, where policy language is silent on the allocation of AI-related risk, insurers may seek to argue than an insured failed to take reasonable precautions to prevent foreseeable losses or, in some cases, breached its duty of utmost good faith.

In such circumstances, questions may arise as to whether a loss resulted from a covered operational failure, the provision of professional service, a regulatory investigation or an uninsured technology-related risk. As these exposures continue to evolve, insurers are increasingly scrutinising AI governance frameworks, model validation processes and oversight mechanisms when underwriting FI risks and assessing claims.

Class actions

As companies and individuals increasingly rely on AI, the risk of AI-related class actions is also increasing. We have already seen the commencement of class actions in the US against AI companies concerning alleged copyright infringement but there is potential for class actions well beyond this arising from the use of AI, and for this to extend to the Australian market. This is consistent with a trend, in the last few years, of class actions being commenced to test the viability of the model in areas of emerging risk, including cyber security, privacy, climate change, and employment.

AI is now being used across the Australian market, in areas including recruitment, financial and professional services, customer engagement, healthcare and digital platforms. However, effective controls are not necessarily applied consistently, increasing the risk of breaches of customer agreements, privacy laws, and other customer-focused regulatory obligations. A single AI-related failure could impact large groups of customers, employees, investors, or consumers causing widespread harm potentially across multiple industries. Beyond direct financial losses, class actions may arise from algorithmic discrimination, privacy breaches, misuse of personal information, AI-driven employment decisions, copyright and training-data disputes, or alleged misrepresentations regarding AI capabilities (i.e., 'AI washing').

In the insurance context, there is scope for coverage 'creep' when claims arise from the use of AI. Such claims may trigger multiple policies, including D&O, PI, Cyber and FI policies, particularly where legacy policy wording does not clearly address AI-related liabilities. As AI risks become more apparent, we expect that organisations with inadequate AI governance, disclosure controls and risk management frameworks are likely to face increased underwriting scrutiny similar to that experienced with the rise of cyber risk.

We're here to help

For more information please contact a member of our team.



David Kerwin
Partner, Directors & Officers
+61 7 3016 5128
David.Kerwin@sparke.com.au



Patrick Tuohey
Partner, Professional Indemnity
+61 3 9291 2243
Patrick.Tuohey@sparke.com.au



Dino Liistro
Partner, Financial Institutions
+61 2 9373 3541
Dino.Liistro@sparke.com.au



Jon Tyne
Partner, Class Actions
+61 2 9260 2683
Jonathan.Tyne@sparke.com.au