

3 Ways Policyholders Can Challenge AI Claims Handling

By **Jillian Raines and Jason Gallant** (August 14, 2026)

Discourse surrounding enterprise artificial intelligence often juxtaposes the corporate risks created by AI's error rate with the scale of efficiencies created by more streamlined business functions.

For many companies, this risk calculus feels tolerable thanks perhaps in part to evolving insurance products purporting to provide protection against business risks created by the deployment of AI. Yet, this discourse rarely focuses on the fact that insurance companies themselves are integrating AI and, consequently, are exposed to the same efficiency-optimizing trade-offs as their policyholders.

The trouble is that when insurers gamble with AI and lose, it is their policyholders who likely suffer the consequences in the form of denied claims and drawn-out litigation.

Just in July, Covenir published its 2026 Insurance Operations Leaders Trends Report, which surveyed 152 U.S.-based insurance decision-makers between February and March 2026 on industry trends.[1]

The report found that 70% of surveyed insurance leaders say that "AI is in production, running in live operational functions," and 57% reported that they run AI in "more than one function" within the insurance industry.[2] These leaders noted that the AI's most measurable impacts included "accelerated claims resolution" and "reduced handling time." [3]

But what exactly does this mean for policyholders? While faster insurance payouts sound great on paper, what happens when there are claim disputes?

In the pre-AI era, the claims handling process, despite its flaws, centered accountability: Underwriters document the policy, insureds give timely notice of claims, and claims handlers evaluate and defend coverage decisions. Individual claims handlers are tasked with reviewing policyholders' claims against the insurance contract, assessing whether the claims fall within that coverage, and then communicating that decision to their policyholders and via their written claims files.

When policyholders disagree and litigation ensues, parties seek discovery about the circumstances surrounding the claim denial, and claims handlers are held accountable for their decisions through depositions and, if necessary, with trial testimony. And, in special cases, where an insurer wrongfully denies coverage, a policyholder may have rights to recover against the insurer under a theory of bad faith.

In this sense, the insurance industry is especially vulnerable to the rapid deployment of AI. Yet, insurers are already piloting it in claim intake, coverage analysis and payment optimization, with varying levels of human oversight.

AI in the insurance industry now operates on a spectrum from a decision-support tool on the one hand to full displacement of human labor on the other. But again, what happens



Jillian Raines



Jason Gallant

when AI begins exercising the same judgment without the same personal accountability?

We see three unique challenges — and potential opportunities for policyholders — that are likely to emerge in future coverage litigation as insurers begin to scale AI.

First, AI training data and model design information will likely be discoverable, potentially unlocking new industrywide insights or areas ripe for reform.

Second, insurers' deployment of, or reliance on, payout-optimizing AI software may strengthen policyholders' bad faith claims against their insurers for wrongful claim denials.

And third, the parameters under which an insurer deploys AI software to assist in coverage decisions may itself run afoul of insurers' good faith obligations under the policies they issue.

1. Policyholders should pursue AI training data.

Industry-specific discovery requests routinely include claims and underwriting notes and manuals, coverage memoranda, reserve and reinsurance information, and even an insurer's own litigation guidelines.

From a policyholder's perspective, the overall theme of this discovery is to vet the basis for an insurer's coverage decision — often exposing flaws or mistakes in the adjustment process. Where AI is integrated into that coverage decision, the data underlying that integration should be fair game for discovery, too.

And to be sure, AI models are characterized just as much by their inputs as they are by their outputs: Flawed reasoning is tainted by flawed front-end training and prompting. Thus, when seeking AI-related discovery from their insurers, policyholders going forward should draw from whatever new categories of information upon which an insurer's AI model rely.

As a first step, policyholders can seek to discover training data ingested by the AI model, including training datasets, model documentation and testing reports. When claims data is ingested into the model, policyholders will want logs of all prompts entered in relation to their claim, as well as any override logs, audit trails or version histories.

And more broadly, policyholders will want to understand the prime directive of an insurer's software — particularly if it is to minimize payout — by obtaining an insurer's AI governance policies.

2. Insurers are at increased risk of bad faith claims.

Beyond obtaining the data to understand how an insurer uses AI, what about the implications of AI usage for an insurer's duty to resolve claims in good faith? Can an insurer satisfy its duty when material aspects of the claims process are delegated to an algorithm — especially if that algorithm is designed to minimize payouts?

Every insurance policy is a contract, and every contract binds the insurer to the implied covenant of good faith and fair dealing. This requires the insurer to pursue an honest and reasonable investigation and evaluation of claims. To be sure, that duty to act in good faith extends to the tools the insurance company deploys to assist that work.

If AI is siloed to just decision-support functions, the adjuster remains ultimately responsible

for claims decisions regardless; a human will ultimately answer for tech failures on dispositive issues. But if AI is authorized to make decisions equivalent to those of an adjuster, insurers must be prepared to take full responsibility for the output.

Marrying payout-minimizing software with merits-based claim evaluations would seem to create serious questions. And the impact of claims decisions that flow from utilizing software potentially biased against coverage compounds when considering that the AI models train on their own payout-minimizing decisions.

Thus, evaluating whether an insurer has carried out its obligation to act in good faith and fair dealing may take on new considerations when accounting for AI-driven errors being widespread — and the fact that AI models often train on those faulty decisions.

Policyholders should seek transparency into how insurers deploy and train their AI tools. While depositions of records custodians are common, so, too, may become AI custodial dispositions — affording policyholders an opportunity to vet how their insurers integrate AI and to what extent an insurer's AI model trains against its own data.

While insurers also historically object to discovery into other claims data and resolutions, insurers may have less success shielding such data from discovery where policyholders need to understand what data was used in an insurer's AI-driven merits decision or in training an AI model used to deny that policyholder's claim.

3. AI claims handling may breach insurance policies.

Current litigation on the issue of whether AI-assisted claims decisions violate insurers' duties of good faith is limited, largely contained to the health insurance context, without courts ruling one way or the other on this issue.

Still, existing precedents may offer protection to policyholders without the need to craft new or even stronger provisions into insurance policies themselves.

Courts have long disapproved of an insurer's failure to conduct reasonable investigations of claims, or of an insurer's placement of its own interests above its insured's. And ironically, the best evidence that courts will not absolve insurers from liability for AI-induced errors might be the ever-growing list of attorneys disciplined by courts for AI-induced errors in legal filings; ignorance of the error or of the technology is simply no excuse from liability.

Broadly, putative litigants should explore the possibility that their insurers failed to oversee the design, deployment or monitoring of their AI systems. The proof would be systemic claim denials, biased outcomes or regulatory violations; the legal claims would be a failure to cooperate and potentially bad faith.

Insurers have the data to show why adverse outcomes occurred in the first place, and the industry will need to be prepared to be held to account if they integrate faulty AI output into adverse claims determinations.

Conclusion

AI will become increasingly embedded in claims handling. Efficiency, however, is not a defense to unreasonable or opaque claims practices, and the coverage litigation landscape will surely adapt: Parties will seek more AI-focused discovery; bad faith doctrines will be tested against both algorithmic and human decision-makers alike; and insurers' good faith

obligations will likely evolve to ensure policyholders remain protected where insurers use AI models instead of adjusters to resolve disputed claims.

Policyholders who remain diligent in discovery and keep abreast of the industry's adoption of AI will be well equipped to make sure they hold their insurers to task in providing the coverage for which they paid.

Jillian Raines is a partner and Jason Gallant is an associate at Cohen Ziffer Frenchman & McKenna.

The opinions expressed are those of the author(s) and do not necessarily reflect the views of their employer, its clients, or Portfolio Media Inc., or any of its or their respective affiliates. This article is for general information purposes and is not intended to be and should not be taken as legal advice.

[1] Covenir, 2026 Insurance Operations Leaders Trends Report: Optimistic, Under Pressure, and Evolving Fast (2026).

[2] Id. at 11.

[3] Id. at 12.