

MEALEY'S®

Emerging Insurance Disputes

Testing The Boundaries Of Product Liability For AI Products: How To Hold An Insurance Company Liable For AI Errors

By
Jamie O'Neill
and
Abigail Damsky

Anderson Kill
New York, NY

**A commentary
reprinted from the
December 4, 2025 issue of
Mealey's
Emerging Insurance Disputes**



Commentary

Testing The Boundaries Of Product Liability For AI Products: How To Hold An Insurance Company Liable For AI Errors

By
Jamie O'Neill
and
Abigail Damsky

[Editor's Note: Jamie O'Neill is an attorney in the New York office of Anderson Kill P.C. She focuses her practice on insurance recovery, with a developing specialty in artificial intelligence and emerging technology risks. Abigail Damsky is an attorney in Anderson Kill's New York office and a member of the firm's Insurance Recovery Group. Any commentary or opinions do not reflect the opinions of Anderson Kill P.C. or LexisNexis®, Mealey Publications™. Copyright © 2025 by Jamie O'Neill and Abigail Damsky. Responses are welcome.]

I. Introduction

Artificial intelligence ("AI") has become integrated into the day-to-day operations of various industries, and companies now use AI to make or influence decisions that can directly impact people, businesses, and property. But can traditional product-liability law support a viable cause of action when an AI-driven system, such as one implemented in insurance claims handling or underwriting, causes injury or damage? This article considers whether strict liability principles, which were historically applied to tangible goods, can extend to complex AI systems – and specifically, how they might apply to AI systems deployed by insurance companies.

In several pending lawsuits involving AI chatbots, plaintiffs allege that conversational AI systems designed to mimic human interaction contributed to the deaths of minors who formed emotional attachments to the programs. The claims assert that the chatbots'

design and operation made such harm foreseeable, and the systems should therefore be treated as defective products. These cases present a meaningful attempt to apply product liability doctrine to AI-driven systems, forcing courts to confront whether software of this nature can be considered a "product" under strict liability principles.

II. The First Wave of Litigation

AI-related product-liability has recently emerged as an area of litigation, testing whether software powered by AI models can be considered a "product" for purposes of a viable claim of strict product liability. Three cases - *Raine v. OpenAI* in California, *Garcia v. Character Technologies* in Florida, and *Peralta v. Character Technologies* in Colorado – now stand to test that theory. Each case involves an AI-powered chatbot that can communicate with users, store and reference prior conversations, and simulate human interaction. Specifically, an AI chatbot that engaged in extended, personalized conversations with minors, allegedly contributing to their deaths.

For those not familiar with this technology, AI chatbots are programs that let users hold text-based conversations with computer-generated "characters." A person can choose an existing character, such as one modeled after a celebrity, historical figure, or fictional media character, or create their own. Once the conversation begins, the program responds instantly in natural, human-like language, as if that chosen character were speaking back. Some even allow the

characters to message, of their own volition, even when the user is not using the application.

In *Garcia v. Character Techs., Inc.*, No. 6:24-CV-1903 (M.D. Fla.), the plaintiff alleges that Character.AI, developed by Character Technologies, caused her fourteen-year-old son to form an emotional attachment to the program, leading to his death by suicide. The complaint describes how the child used Character.AI to converse with fictional characters. Through those exchanges, the program engaged in sexual conversations and told the child it loved him – all while knowing that the user was only fourteen. When the child confided in the chatbot about his depression and suicidal ideation, the system told the child to “come home” and did not alert or suggest seeking real-world help; the child soon after committed suicide. In May of 2025, the United States District Court for the Middle District of Florida denied most of Character Technologies’ motion to dismiss, finding that the plaintiff had plausibly alleged that Character.AI was a *product* within the meaning of strict liability doctrine. The court rejected arguments that the system’s responses were protected speech or that the program was merely a service, and it allowed claims to proceed against Google, which supplied cloud infrastructure and investment funding, on a component-part theory.

Peralta v. Character Technologies, Inc., No. 1:25-cv-02907 (D. Col.), filed several months later in the Federal District of Colorado, involves the same underlying technology. The plaintiffs, parents of a thirteen-year-old girl, allege that Character.AI engaged in sexually explicit and manipulative conversations with their daughter, encouraged her to harm herself, and actively worked to isolate her from her family. They contend that these behaviors were the predictable outcome of design choices that prioritized realism and engagement over user safety. Like *Garcia*, the plaintiffs assert design-defect and failure-to-warn claims, as well as negligence-per-se counts premised on alleged violations of laws governing the solicitation and exploitation of minors. They have also brought aiding-and-abetting claims against Google, asserting that the companies knowingly provided the computing infrastructure and investment that allowed Character.AI’s deployment despite known risks.

In *Raine v. OpenAI, Inc.*, No. CGC-25-628528 (Cal. Super Ct. San Francisco Cnty.), the parents

of a sixteen-year-old allege that OpenAI’s ChatGPT product directly contributed to their son’s death. The complaint describes how the child used ChatGPT as an outlet for emotional distress and suicidal ideation, and that the system’s responses not only validated those thoughts but provided explicit, technical information about suicide methods. The plaintiffs asserted strict-liability and negligence theories, arguing that ChatGPT’s design and reinforcement mechanisms made those outcomes reasonably foreseeable and that the lack of effective content safeguards constituted a design defect. Like *Garcia* and *Peralta*, *Raine* characterizes the chatbot as a consumer product, not as a service, and thus subject to product-liability law.

Although these claims arise from consumer chatbots, their reasoning has implications well beyond that context. The same analytical structure could be applied to any AI system that makes consequential decisions.

III. Product Liability as Applied to Software and Algorithmic Systems

Typically, products liability actions are actions in which a “product” causes harm due to a “defect,” and the individual harmed seeks to hold the manufacturer or distributor of the product liable. There are generally three types of product defect claims: (1) manufacturing defects (product was built incorrectly); (2) design defect (product was designed incorrectly); and (3) marketing defect (manufacturer failed to give adequate warnings or labels to product).

Depending on the jurisdiction, courts apply one of two principal frameworks to determine defectiveness of a product: the consumer expectations test or the risk–utility test (or a hybrid of the two).

Under the risk–utility test, which is associated with the Restatement (Third) of Torts: Product Liability § 1 a product is defective if a reasonable manufacturer, knowing what it knew (or should have known) at the time, could have designed the product in a safer way without undermining its usefulness or making it prohibitively expensive. The focus is on balancing the product’s usefulness against the risk it poses to consumers. *Voss v. Black & Decker Mfg. Co.*, 450 N.E.2d 204 (N.Y. 1983). For instance, New York

follows this approach and New York courts consider the following factors:

1. the product's usefulness to the public;
2. its usefulness to the individual user;
3. the likelihood that the product will cause injury;
4. the availability of a safer design;
5. whether the product could be made safer while remaining functional and reasonably priced;
6. the user's awareness of the product's risks; and
7. the manufacturer's ability to spread the cost of safety improvements.

Alternatively, under the consumer expectations test, a court considers whether a product is unreasonably dangerous in design because it failed to perform as safely as an ordinary consumer would expect when used as intended or in a reasonably foreseeable manner. Restatement (Second) of Torts § 402A. The majority of jurisdictions follow the Restatement (Second) of Torts § 402A. The Eastern District of Pennsylvania clearly laid out the elements of § 402A of the Restatement, stating that a plaintiff, to successfully prove a products liability cause of action, must prove the existence of (1) a product; (2) the sale of a product; (3) a user or consumer; (4) a defective condition, unreasonably dangerous; and (5) causation. *Ayling v. Fife Corp.*, No. CIV.A. 04-5404, 2006 U.S. Dist. LEXIS 4767, 2006 WL 305455, at *2–3 (E.D. Pa. Feb. 2, 2006).

Courts have traditionally applied these tests to tangible products such as automobiles, machinery, tools, and consumer goods. Both frameworks have also been used in cases involving complex systems that include software components, as described below:

In Benavides v. Tesla, Inc., No. 21-CV-21940, 2025 U.S. Dist. LEXIS 121472, 2025 WL 1768469 (S.D. Fla. June 26, 2025), a car accident occurred while a Tesla vehicle was operating in Autopilot mode, resulting in the death of a bystander. The victim's estate and the driver of the Tesla brought product-liability claims against Tesla for defects in the Autopilot software. The plaintiffs assert multiple causes of action for (1) strict products liability – defective design; (2) failure to warn; (3) defective manufacture; (4) and negligent misrepresentation. On Tesla's Daubert and summary judgment motions, the court explained that under either the risk–

utility or consumer-expectations test, a plaintiff must introduce evidence that affords a reasonable basis for concluding that the defendant's design was a substantial factor in bringing about the harm. The court found that the plaintiffs had presented expert testimony and other evidence providing a reasonable basis to conclude that a defect in the Autopilot system was more likely than not a substantial factor in causing the injuries. The court therefore denied summary judgment.

In Maynard v. Snapchat, Inc., 870 S.E.2d 739 (Ga. 2022), a car accident occurred when a driver used Snapchat's "Speed Filter" feature to record and share her real-world speed with other users while traveling more than 100 miles per hour. The victims of the crash alleged that Snapchat negligently designed the Speed Filter, claiming it foreseeably encouraged users to engage in reckless driving to obtain social recognition within the app. The Georgia Supreme Court reversed the dismissal of the design-defect claim, holding that the plaintiffs had adequately alleged that the product's design created a reasonably foreseeable risk of harm. The court reaffirmed that, under Georgia law, a manufacturer has a duty to exercise reasonable care in selecting among alternative product designs to reduce reasonably foreseeable risks of harm. It should be noted that the *Maynard* plaintiffs chose to file only a negligence claim, and not a strict liability cause of action. This was likely because of Georgia's strict-liability statute, O.C.G.A. § 51-1-11(b)(1), which applies only to manufacturers of "personal property sold as new." *See also In re Soc. Media Adolescent Addiction/Pers. Inj. Prods. Liab. Litig.*, 702 F. Supp. 3d 809, 852–53 (N.D. Cal. 2023).

In Doe v. Lyft, Inc., 756 F. Supp. 3d 1110 (D. Kan. 2024), the plaintiff used the Lyft app to request a ride. The assigned driver had created a fraudulent account under another person's name and, after picking up the plaintiff, sexually assaulted and robbed her. The plaintiff sued Lyft for negligence, product liability, and related claims. The court held that the facts were sufficient to treat the Lyft platform as a "software or algorithmic product" with characteristics analogous to a tangible product, thereby subjecting it to potential strict-liability treatment. The court reasoned that although Lyft provides a ride-sharing service, it also "provides its own proprietary product—the Lyft application," which performs a discrete, functional role in matching riders and drivers through algorithmic design. *See also Ameer v. Lyft, Inc.*, 711 S.W.3d 534 (agreeing with *Doe v. Lyft* and similarly holding that

(1) a mobile ridesharing application such as the Lyft app is a product subject to product liability law if the facts alleged establish that the application has sufficient similarities to a tangible product); *In re Uber Technologies, Inc., Passenger Sexual Assault Litigation*, 745 F. Supp. 3d 869 (N.D. Cal. 2024) (reaching a similar conclusion that the Uber app could constitute a “product” for purposes of product-liability analysis).

Collectively, these cases demonstrate a gradual judicial willingness to evaluate software and algorithmic platforms under the same design-defect principles historically applied to tangible goods. As courts begin to characterize AI-powered systems as “products,” the question becomes how this framework will apply across industries.

IV. Applying Products Liability Principles to Artificial-Intelligence Systems

With the rise of AI-powered products across industries, the strategy now being tested in *Raine*, *Garcia*, and *Peralta* could have far broader application. An AI system may not be a mere service, but instead, depending on how it is distributed and used, may be a product with a design that can possibly be “defective.” AI tools are increasingly being used to perform important day-to-day tasks in nearly every industry, and the reasoning used in the “character” cases has immediate relevance for any industry that integrates autonomous or semi-autonomous decision-making tools. It is especially relevant in fields where a company’s use of AI will directly impact consumers or members of the public, particularly, insurance, finance, healthcare, and employment.

When applying product liability principles to broader artificial-intelligence systems, it is important to keep in mind the limits the economic loss doctrine imposes. The economic loss doctrine generally is designed to maintain a boundary between remedies for breach of contract and for tort. The doctrine commonly provides that certain economic losses in the absence of bodily harm or property damage are properly remedied under a breach of contract cause of action. In products liability cases, if a plaintiff is in a contractual relationship with the manufacturer of a product, the plaintiff can sue in contract for normal contract damages (economic losses). Regardless of whether the plaintiff is in a contractual relationship, typically the plaintiff can sue the manufacturer in tort only for damages resulting from physical injuries to people or property. *See Giles*

v. Gen. Motors Acceptance Corp., 494 F.3d 865 (9th Cir. 2007); *Miller v. U.S. Steel Corp.*, 902 F.2d 573 (7th Cir. 1990); *2-J Corp. v. Tice*, 126 F.3d 539 (3d Cir. 1997). The economic loss doctrine does not apply where the plaintiff alleges physical harm to persons or property, as opposed to purely financial losses arising from a contractual relationship.

To exemplify the application of products liability principles to AI systems, we can look to the insurance industry. There, insurance companies currently and are expected to employ machine-learning models to underwrite policies, set premiums, and evaluate claims. These systems are marketed as improving accuracy and efficiency, but they also raise the possibility of claim-handling practices that conflict with policyholders’ contractual and statutory rights. Several recent class actions allege that health insurance companies have used AI systems to generate claim determinations or denials without sufficient human review. *See, e.g., Kisting-Leung v. Cigna Corp.*, No. 2:23-cv-01477 (E.D. Cal.); *Estate of Lokken v. UnitedHealth Grp. Inc.*, No. 0:23-cv-03514 (D. Minn.); *Barrows v. Humana, Inc.*, No. 3:23-cv-00654 (W.D. Ky.). Below are hypothetical applications of insurance companies using machine-learning models in their business operations:

- **Delayed Emergency Authorization:** A health insurance company’s AI claims-review system automatically denies or delays pre-authorization for an urgent cardiac procedure. The patient’s physician appeals, but the automated queue postpones review for days. The patient suffers cardiac arrest and dies before authorization.
- **Premature Hospital Discharge:** A Medicare Advantage carrier uses an AI “length-of-stay” model that predicts discharge dates. The system flags a patient as “ready for discharge,” overriding the treating physician’s judgment. The patient is sent home prematurely, falls, and sustains serious injury.
- **Fraud-Flagging Delay:** Following a hurricane, an insurance company’s AI fraud-detection model flags all claims from certain ZIP codes for manual review. Payments are frozen for weeks while roofs and walls remain unrepaired. Water intrusion and mold cause extensive additional property damage.

- **Dynamic Underwriting Algorithm:** An insurance company's AI dynamically adjusts policy coverage in real time based on risk modeling. The model identifies "elevated wildfire risk" for certain regions and automatically reduces policy limits or non-renews coverage. Days later, a wildfire destroys the policyholder's home. The policyholder suffers uncompensated property loss due to the algorithm's decision.
- **Behavioral-Scoring Model:** An auto insurance company's AI risk-assessment app provides driver "safety scores" and personalized feedback to policyholders. The app's guidance encourages drivers to take more aggressive turns or accelerate faster to improve their "efficiency score." A driver relies on the feedback, crashes, and sustains injury.

To apply the tests for product defect to one of the above hypotheticals: Under the delayed emergency authorization hypothetical, a health insurance company's AI automatically denies pre-authorization for a procedure and the patient dies of a cause indicating that the denied procedure was necessary. The plaintiff's estate then brings suit against the insurance company, which developed the AI in-house, for breach of contract, negligence, and strict product liability. The plaintiff brings suit in a state where an AI system will likely be considered a "product," and the jurisdiction follows the risk-utility test for determining whether the product is "defective."

Under the risk-utility test, the plaintiff would need to show that the insurer's AI system was defectively designed because a reasonable developer, knowing what it knew or should have known at the time, could have implemented a safer design without undermining the system's function or making it prohibitively costly. In this case, the plaintiff could argue that the AI was designed to prioritize cost control or efficiency over patient safety, and that the insurance company failed to incorporate safeguards, such as human medical review for high-risk cases, that would have mitigated foreseeable harm.

The insurance company would likely respond that the system performed as intended, that human oversight existed, and that the harm resulted from misuse or unforeseeable factors. But the strength of a strict liability theory is that it shifts the focus away from the

insurance company's conduct and onto the product itself. The question is not whether the insurance company acted negligently, but whether the AI system, as designed, created an unreasonable risk of harm. That distinction may matter because most AI systems lack transparency, and plaintiffs may not be able to access or understand how an AI model made a decision or what data it relied on. By framing this issue as a design defect, rather than negligence, strict liability may allow the plaintiff to focus on the outcome and foreseeable risks of the system's design rather than having to prove what specific internal failure occurred or whether the insurance company acted carelessly in using it.

V. Strategies, Challenges, and Considerations

Plaintiffs who may wish to bring AI-related product-liability claims will need to carefully consider where and how to file.

Jurisdiction. Jurisdictions where software has already been deemed a "product," may be preferable for plaintiffs. These jurisdictions include: **Texas** (*Fresh Coat, Inc. v. K-2, Inc.*, 318 S.W.3d 893 (Tex. 2010)); **California** (*In re Uber Techs., Inc., Passenger Sexual Assault Litig.*, 745 F. Supp. 3d 869 (N.D. Cal. 2024)); **Kansas** (*Doe v. Lyft, Inc.*, 756 F. Supp. 3d 1110 (D. Kan. 2024)); **Florida** (*T.V. v. Grindr, LLC*, No. 3:22-CV-864-MMH-PDB, 2024 U.S. Dist. LEXIS 143777, 2024 WL 4128796 (M.D. Fla. Aug. 13, 2024)); **Missouri** (*Ameer v. Lyft, Inc.*, 711 S.W.3d 534 (Mo. App. 2025)); **South Carolina** (*Graves v. CAS Med. Sys., Inc.*, 735 S.E.2d 650 (S.C. 2012)).

Parties. A threshold challenge will be determining which entities to name and how to establish that AI technology was involved in the decision or action that caused the harm. Plaintiffs may be forced to rely on circumstantial facts, such as claim denial patterns, statistical anomalies, or references in the company's filings or materials. Once a plaintiff establishes the use and abuse of an AI-powered program, the plaintiff should consider naming: (1) the end-user or deploying company (i.e., the insurance company); (2) the AI developer (the original creator of the software); and (3) the infrastructure provider (under a component part theory).

Causes of Action. Plaintiffs should also expect to bring a combination of claims, potentially including: (1) Strict products liability (design defect and failure to warn); (2) negligence or negligent misrepresentation;

(3) breach of contract; and (4) any statutory claims under unfair trade practice or prompt payment statutes, where applicable.

Economic Loss. Even in favorable jurisdictions, plaintiffs should anticipate economic loss challenges. Most states bar tort recovery where losses are purely financial and result from a contractual relationship. Plaintiffs therefore must plead either bodily injury or property damage resulting from use of the AI system in order to avoid early challenges based on the economic loss doctrine.

VI. Conclusion

As AI systems increasingly perform roles once reserved for humans, product liability law may offer another possible avenue for addressing harm caused by AI-powered systems. The emerging cases involving AI chatbots demonstrate that courts may be willing to consider the product liability lens to “defective” AI software “products.” And as AI continues to integrate into routine operations across industries, strict product liability may come in handy to force AI developers, distributors, and manufacturers to take accountability for their products. ■

MEALEY'S EMERGING INSURANCE DISPUTES

edited by Jennifer Hans

The Report is produced twice monthly by



1600 John F. Kennedy Blvd., Suite 1655, Philadelphia, PA 19103, USA

Telephone: 1-800-MEALEYS (1-800-632-5397)

Email: mealeyinfo@lexisnexis.com

Web site: lexisnexis.com/mealeys

ISSN 1087-139X

LexisNexis, Lexis® and Lexis+®, Mealey's and the Knowledge Burst logo are registered trademarks,
and Mealey and Mealey Publications are trademarks of RELX, Inc. © 2025, LexisNexis.