
FEBRUARY RESPONSES

THE QUESTIONABLE ASSUMPTIONS UNDERPINNING LIABILITY-INSURANCE PROPOSALS[†]

KEVIN FRAZIER^{*}

CONTENTS

INTRODUCTION	1262
I. ASSUMPTION I: PUBLIC POLICY SUPPORTS UPENDING PRECEDENT	1262
II. ASSUMPTION II: EXPANSION OF LIABILITY WILL BENEFIT THOSE HARMED BY AI	1264
III. ASSUMPTION III: LITIGATION WILL RESULT IN THE “RIGHT KIND” OF INNOVATION	1266
CONCLUSION	1267

[†] An invited response to Renee Henson, *Artificial Intelligence, Judicial Evolution, and Insurance*, 106 B.U. L. REV. 241 (2026).

^{*} AI Innovation and Law Fellow, The University of Texas School of Law. Sven Lohse provided excellent research assistance in the formation of this Response.

INTRODUCTION

Professor Renee Henson makes a methodical, logical, and, for the most part, descriptive case for the emergence of artificial intelligence (“AI”) liability insurance.¹ Shifting jurisprudence on product liability, Section 230 of the Communications Decency Act, and the First Amendment carries the potential to expose AI actors—large and small—to substantial liability.² Existing insurance coverage, however, likely does not cover such cases, leaving innovators in a precarious financial position—the sort of financial risk that threatens to chill innovative activity in the first place.³ Professor Henson anticipates that new forms of insurance will emerge to fill that void.⁴ None of this is controversial, and perhaps that’s the problem.

Professor Henson’s underlying theory, which sits atop a mountain of assumptions that have hardened into shibboleths, is as follows: (1) Public policy concerns justify judicial reinterpretation of existing law to extend liability to cover emerging issues;⁵ (2) the extension of liability is met by an expansion in insurance coverage to allow for small and large AI companies to avoid financial collapse⁶ and for those negatively impacted by AI to recover damages;⁷ and (3) as a result some acceptable amount of innovation will come about.⁸ More specifically, Professor Henson explicitly calls for expansive third-party algorithmic liability insurance terms, envisioning policies with broad definitions of injury and occurrence and few exclusions.⁹

However, further analysis of each of these three assumptions flips Professor Henson’s conclusion on its head. Rather than produce the societally optimal level of innovation, her theory risks the creation of a liability-insurance-anti-innovation loop. Professor Henson’s summation of the facts of recent litigation shows how this is already playing out in practice.

I. ASSUMPTION I: PUBLIC POLICY SUPPORTS UPENDING PRECEDENT

Courts have succumbed to public pressure to adopt creative interpretations of case law to open the door to liability. Professor Henson details how courts have excused their departures from existing law, by in part, relying on public policy concerns. For example, in the case *In re Social Media Adolescent Ad-*

¹ Renee Henson, *Artificial Intelligence, Judicial Evolution, and Insurance*, 106 B.U. L. REV. 241 (2026).

² *Id.* at 244–49; *id.* at 284 n.255.

³ *Id.* at 250–53.

⁴ *Id.* at 248–50; *id.* at 248 n.27.

⁵ *Id.* at 265; *see id.* at 266 n.124.

⁶ *Id.* at 253.

⁷ *Id.* at 248–50.

⁸ *Id.* at 299–301.

⁹ *Id.* Section IV.B.

diction/Personal Injury Products Liability Litigation,¹⁰ a matter involving the use of social media platforms by minors, the court reinforced its expanded sense of what qualifies as a “product” in product liability suits by pointing to the Third Restatement’s recommendation to weigh “the public’s interest in life and health.”¹¹ A cursory policy analysis by the court contended that mitigating the alleged defects with the apps was in the public interest. The court flagged that the well-being of children is at stake, that society cares about kids, and, therefore, determined that an expanded interpretation of “product” to allow such suits would “advance the public’s interest in young people’s well-being.”¹²

Yet, research by the American Psychological Association (“APA”) reveals a more nuanced policy environment.¹³ Rather than lump minor users into one category, the APA notes that the benefits and costs of social media use change depending on the precise age range of the child.¹⁴ Whereas social media use between the ages of eleven to thirteen for girls and fourteen to fifteen for boys may warrant heightened scrutiny, “[a]s kids become more mature and develop digital literacy skills, they can earn more autonomy.”¹⁵ Whether social media is a net positive also varies on a child’s mental well-being. The following positive use cases are omitted from the court’s analysis (and Professor Henson’s):

Online social interaction can promote healthy socialization among teens, especially when they’re experiencing stress or social isolation. For youth who have anxiety or struggle in social situations, practicing conversations over social media can be an important step toward feeling more comfortable interacting with peers in person. Social media can also help kids stay in touch with their support networks. That can be especially important for kids from marginalized groups, such as LGBTQ+ adolescents who may be reluctant or unable to discuss their identity with caregivers.¹⁶

Had the court been made aware of these policy nuances, it may have been less likely to push the bounds of precedent.

Courts encountering lawsuits arising from AI use stand on even shakier ground to the extent they think public policy aligns with the expansion of

¹⁰ 702 F. Supp. 3d 809 (N.D. Cal. 2023).

¹¹ *Id.* at 850 n.50 (citing RESTATEMENT (THIRD) OF TORTS § 19 cmt. a (A.L.I. 1998)).

¹² *Id.*

¹³ See Kirsten Weir, *Protecting Teens on Social Media*, 54 MONITOR ON PSYCH. 46, 48 (2023) (acknowledging benefits of social media use by teens).

¹⁴ See *id.* at 49-51.

¹⁵ *Id.* at 49-50.

¹⁶ *Id.* at 50 (citing Shelley L. Craig, Andrew D. Eaton, Lauren B. McInroy, Vivian W. Y. Leung & Sreedevi Krishnan, *Can Social Media Participation Enhance LGBTQ+ Youth Well-Being? Development of the Social Media Benefits Scale*, 7 SOC. MEDIA + SOC’Y, Jan.-Mar. 2021, at 1).

existing law. A survey of ChatGPT users revealed that only 1.9% of all use cases are for “relationships”¹⁷—the type of use that has elicited headlines and resulted in lawsuits.¹⁸ It’s unknown how many of those uses are by minor users nor how many of those uses are problematic. For courts to assume that the policy scales weigh in favor of limiting or even banning the use of AI tools by minors to form some sort of relationship is premature.

To the extent that courts willing to adopt new interpretations of the law are instead motivated by a desire to shield adults from such relationships, the policy picture is even murkier. In a system such as ours in which the law is intended to preserve space for individual autonomy and freedom of thought, it is questionable whether courts have the authority to intervene in how adults opt to engage with AI tools.¹⁹

Even beyond liability arising from relationship-based uses of AI, the public policy justification for perpetuating the liability-insurance loop is unclear. The public policy justification is addressed more fully under Assumption III.

II. ASSUMPTION II: EXPANSION OF LIABILITY WILL BENEFIT THOSE HARMED BY AI

The call for the expansion of the liability-insurance framework also rests on the idea that doing so will further the interests of harmed parties and society generally. If the civil justice system worked as it is outlined in textbooks, then this idea would likely be true. However, expeditious litigation based on high-quality information applying clear laws that results in objectively appropriate attorneys’ fees and damages is as common as a unicorn.

A scan of the dollars and cents of personal injury litigation demonstrates this point. By one estimate, the average personal injury jury award was \$125,366 as of 2020.²⁰ By another estimate, the average product liability settlement lands in the \$30,000 to \$100,000 range.²¹ Pursuit of defective product

¹⁷ Aaron Chatterji et al., *How People Use ChatGPT 2* (Nat’l Bureau of Econ. Rsch., Working Paper No. 34255, 2025).

¹⁸ See, e.g., Henson, *supra* note 1, at 244 (discussing lawsuit by parents after child “developed a romantic, addictive relationship with [a] chatbot, which allegedly led to mental and emotional distress and ultimately suicide”).

¹⁹ See Adam J. Kolber, *Two Views of First Amendment Thought Privacy*, 18 U. PA. J. CONST. L. 1381, 1391 (2016) (discussing long history of constitutional protection for thought privacy).

²⁰ See *Facts + Statistics: Product Liability*, INS. INFO. INST., <https://www.iii.org/fact-statistic/facts-statistics-product-liability> [<https://perma.cc/85FX-LPAA>] (last visited Feb. 2, 2026).

²¹ See Omar Ochoa, *What is the Average Product Liability Settlement in Texas in 2026*, OMAR OCHOA L. FIRM (Nov. 26, 2025), <https://www.omarochoalaw.com/blog/average-product-liability-settlement> [<https://perma.cc/E2LK-DSH2>].

cases commonly generates more than \$100,000 in expenses.²² Assuming that a successful AI-related liability claim necessitates hiring one or more AI experts, it's fair to estimate that those costs will only increase in the cases imagined by Professor Henson. These figures alone warrant at least pausing to consider if doubling down on the current civil litigation system is a wise strategy for responding to the very real need to protect and make whole consumers injured by AI-related incidents. Accounting for the social costs of such litigation emphasizes the need for reevaluation.

Research by Professor Jay Tidmarsh suggests that the math behind the social cost-benefit calculation of such litigation is damning for those contending that litigation is the best or even an appropriate path forward.²³ He defines social benefits as “the combined benefits that litigation provides to the private individuals involved in the litigation and to society.”²⁴ Proponents of a litigation-forward approach may claim that it produces positive externalities by way of encouraging AI labs to deploy safer AI tools. Put differently, the thinking goes that the public benefits when litigation demands labs to update their development and deployment practices to reduce the odds of injury.

Social costs refer to “both the costs incurred by the litigants as well as the costs imposed on society.”²⁵ Such costs may include court expenses, which, for the sake of establishing a baseline, exceed \$10,000 per day in California.²⁶ Given that the private costs discussed above already indicate that this assumption is not valid, inclusion of the net social benefit only bolsters the conclusion that litigation will not serve harmed parties and the general public. From a social benefits perspective, it's unlikely that litigation is the only or more efficient way of facilitating improvements to AI tools. Myriad other pressures can drive consumer-friendly innovation, including but not limited to negative press, consumer demands, industry standards, guidance documents from regulators, and competition from rivals and upstarts.²⁷ From a social costs perspective, the \$10,000 per day estimate in California likely understates the social costs of AI-liability cases due to the expenses of the rele-

²² See *How Much Will It Cost Me to Hire a Defective Products Lawyer*, VILES & BECKMAN, LLC, <https://www.vilesandbeckman.com/faq/how-much-will-it-cost-to-hire-defective-product-lawyer> [<https://perma.cc/4BL6-M8N4>] (last visited Feb. 2, 2026).

²³ See Jay Tidmarsh, *The Litigation Budget*, 68 VAND. L. REV. 855, 877, 881-82 (2015); *id.* at 860 (“[L]itigation should not cost more than the benefit it provides to society.”).

²⁴ *Id.* at 861 n.17 (emphasis omitted).

²⁵ *Id.*

²⁶ CAL. SEN. RULES COMM., OFF. OF SEN. FLOOR ANALYSES, THIRD READING OF ASSEMBLY BILL NUMBER, 2025-2026 Reg. Sess., at 1064.

²⁷ See WebFX, *The History of Aviation Safety*, HRD AERO SYSTEMS (June 15, 2022), <https://www.hrd-aerosystems.com/blog/history-of-aviation-safety> [<https://perma.cc/WP56-7TP7>] (explaining evolution of aviation industry was influenced by numerous non-litigation pressure).

vant experts²⁸ and the opportunity costs of tying innovators up in court proceedings while they could be working on the next generation of their AI tool.

III. ASSUMPTION III: LITIGATION WILL RESULT IN THE “RIGHT KIND” OF INNOVATION

In the event that there’s an unequivocal case for an explicit public policy intervention arising from an AI flaw and the math seemingly suggests that litigation is a net social positive, one final assumption still requires investigation: The resulting incentives will steer AI innovation in a socially optimal fashion. Litigation-forward proponents of AI regulations assume that products liability cases form guardrails that direct innovative activity away from AI strategies that courts have deemed off limits. This thinking may hold true in fields pertaining to fairly static and well-understood technologies. For example, few would contest the merits of lawsuits that have accelerated the adoption of rearview cameras by automobile manufacturers.²⁹ How to safely develop AI in a manner that reflects the socially optimal mix of pros and cons is far less clear. Litigation at this early stage in development risks diverting labs away from new research and new models that might unleash significant societal benefits.³⁰

AI experts attest that they do not fully understand how models “learn” nor can they predict with complete accuracy which capabilities models will demonstrate once deployed.³¹ Efforts to evaluate whether models are aligned with user expectations and sufficiently incapable of engaging in certain problematic behavior are at an early stage.³² If litigation has the effect of encouraging labs to spend less time on novel inquiries and to instead prioritize tools

²⁸ *Expert Witness Fees - How Much Should an Expert Witness Charge?*, SEAK EXPERT WITNESS DISCOVERY, <https://blog.seakexperts.com/expert-witness-fees-how-much-should-an-expert-witness-charge> [https://perma.cc/8MSU-KZSE] (last visited Feb. 2, 2026) (explaining median hourly fee for expert testimony is \$500 per hour).

²⁹ See Fred Meier & Chris Woodyard, *Administration Sued Over Backup Camera Delay*, USA TODAY (Sep. 26, 2013, at 08:20 ET), <https://www.usatoday.com/story/money/cars/2013/09/25/backup-cameras-dot-lawsuit-gulbransen-obama-administration/2870819/>.

³⁰ See Dario Amodei, *Machines of Loving Grace*, DARIOAMODEI.COM (Oct. 2024), <https://www.darioamodei.com/essay/machines-of-loving-grace> [https://perma.cc/7JMJ-5TDH] (describing immense possible benefits for human society if AI is properly developed).

³¹ See RISHI BOMMASANI ET AL., THE CALIFORNIA REPORT ON FRONTIER AI POLICY 21 (2025), <https://www.gov.ca.gov/wp-content/uploads/2025/06/June-17-2025---The-California-Report-on-Frontier-AI-Policy.pdf> [https://perma.cc/5L9U-MJQR] (discussing complexity of training AI models and observation of AI models developing deceptive behaviors).

³² See Robert Booth, *Experts Find Flaws in Hundreds of Tests That Check AI Safety and Effectiveness*, GUARDIAN (Nov. 3, 2025, at 19:05 ET), <https://www.theguardian.com/technology/2025/nov/04/experts-find-flaws-hundreds-tests-check-ai-safety-effectiveness>.

that fit within the narrow parameters of whatever courts have deemed safe, then labs may never uncover new tactics and methods to improve model reliability, utility, and safety.

Each “successful” suit is another signal to investors and founders to lean more in one direction over another. While that may be fine in settled fields, the potential for unintended consequences that are a net societal negative are significant in emergent fields like AI. Consider, for example, limitations on models that can generate compelling deepfakes. Courts may wind up imposing rules that severely restrict the development and deployment of models capable of detecting deepfakes. Yet the irony is unavoidable: The only reliable way to train models to spot deepfakes is to expose them to vast numbers of them—hundreds of millions, not dozens.³³

The upshot is that the “kind” of innovation that develops through litigation may reflect a common tendency among jurists (and humans generally) to place less weight on known risks compared to unknown risks.³⁴ Ambiguity aversion colors our ability to rationally assess important trade-offs regarding emerging technology. One common example is the frequency with which patients decline efficacious procedures because they simply do not understand how it works.³⁵ In the context of AI-liability litigation, a default willingness to defer to the status quo and accept known risks is especially problematic. Judges may not meaningfully weigh the possibility that the graver risk to social well-being is blunting AI progress and adoption in a particular domain.

CONCLUSION

Professor Henson’s theory risks creating a governance system optimized not for public benefit, but for legibility to underwriters. When insurers—rather than technologists, consumers, or expert regulators—become the de facto arbiters of “safe” innovation, the “kind” of AI we get will be stunted. This is the true mechanism of the anti-innovation loop: Insurers, to manage their own risk, will write policies that demand their clients (AI labs) avoid novel research with risks that are difficult to quantify.³⁶ This forces innovation into narrow, pre-approved channels that are easily insurable, not necessarily beneficial. Henson’s descriptive case for the emergence of insurance thus becomes a prescriptive warning about its dominance. We risk a future

³³ Jeongsoo Park & Andrew Owens, *Community Forensics: Using Thousands of Generators to Train Fake Image Detectors*, ARXIV (Nov. 6, 2024), <https://arxiv.org/pdf/2411.04125v1> [<https://perma.cc/X4SW-XUCA>].

³⁴ See *Why Do We Prefer Options We Know?*, DECISION LAB, <https://thedecisionlab.com/biases/ambiguity-effect> [<https://perma.cc/JP94-J5RG>] (last visited Feb. 2, 2026).

³⁵ *Id.*

³⁶ See KENNETH S. ABRAHAM & DANIEL SCHWARCZ, *INSURANCE LAW & REGULATION* 9 (8th ed. 2025).

where progress is defined not by what is scientifically possible, but by what can be neatly categorized on an actuarial table.

Ultimately, Professor Henson has provided a valuable, if inadvertent, roadmap of a future we must avoid. The emergence of AI-liability insurance isn't a sign of a healthy market adapting; it's a symptom of a foundational mismatch. We are attempting to govern a dynamic, general-purpose technology with the slow, expensive, and backward-looking tools of tort. This framework is destined to fail, satisfying neither victims (due to high costs) nor the public (by chilling progress). The alternative is to reject the liability-insurance framework as our primary governance mechanism and instead build a nimbler, *ex ante* system—one rooted in expert-informed standards, regulatory sandboxes, and clear safe harbors that reward, rather than punish, the very research needed to make AI safer and more capable.