

RETHINKING ROBOT LIABILITY

ZACHARY HENDERSON

INTRODUCTION	480
I. PROBLEMS, REAL AND IMAGINED	483
<i>A. Why Current-Generation AI Actors Are Unique</i>	485
<i>B. Negligence Can Handle AI Harms</i>	488
1. AI Actors Are Not Susceptible to Negligence Claims . . .	489
2. . . . but the Legal Entities That Deploy Them <i>Can</i> Be Negligent.....	490
3. The Indirect Effects of AI on Negligence and the Reasonable Person Standard.....	493
<i>C. AI Has No Products Liability Immunity</i>	494
II. PROPOSED SOLUTIONS AND THEIR SHORTCOMINGS.....	502
<i>A. Strict Liability</i>	503
1. Normative Shortcomings of Strict Liability for AI Harms	506
2. Doctrinal Shortcomings of Strict Liability for AI Harms	510
<i>B. Other Proposed Solutions</i>	513
III. THE CASE FOR LETTING AI HARMS LIE	515
<i>A. Not All Harms Are Tortious</i>	516
<i>B. Corrective Justice Theory</i>	517
<i>C. Economic Efficiency Theory</i>	518
<i>D. Civil Recourse Theory</i>	523
<i>E. Loss Distribution Theory</i>	525
IV. INSURANCE, RISK ASSUMPTION, AND AI CATASTROPHES	526
<i>A. Insurance Markets, Old and New</i>	527
<i>B. AI Catastrophes: Reprising Negligence</i>	529
CONCLUSION	530

RETHINKING ROBOT LIABILITY

ZACHARY HENDERSON*

Abstract: Today a growing chorus of voices in the tort-law literature sings a siren song. AI harms must be redressed through across-the-board strict liability, they croon—the same standard we apply to dynamite users, lion tamers, and faulty-chainsaw makers. Their rationale is understandable: at first blush, many AI harms appear impervious to other traditional theories of tort. In what sense might a self-driving car be negligent? Or how might one prove that an AI’s neural network—a “black box” filled with billions or trillions of inscrutable numbers—is defective under a products-liability theory? But tempting though strict liability may be, categorically applying it to AI harms would be a mistake, for at least three reasons. First, just like human activities, AI activities are not monolithic. An AI-enabled treadmill does not pose the same risk as an AI demolitions robot. Second, our default tort rules will work far better than the chorus suggests. Negligence and products liability are high flexibility, high context doctrines that have a long history of redressing harms caused by novel activities and technologies. Plus, they aren’t our only tools. Where an AI activity presents the potential for catastrophic harm, tailored, domain-specific legislation will be the right tool for the job. And where AI activities cause non-tortious harms, data-driven insurance markets (old and new) will be well suited to fill most compensatory voids. Third, even if something like categorical strict liability just for AI harms were desirable, no current theory of tort law could justify the bifurcated tort system that kind of rule would create.

INTRODUCTION

At the height of the Electric Wire Panic of 1889, terrified New Yorkers were ripping wires out of their walls.¹ Thomas Edison, in one of his many hit-pieces against alternating current, warned residents against unlocking their

* Incoming Assistant Professor of Law, SMU Dedman School of Law; Visiting Assistant Professor and Shashi and Dipanjan “DJ” Deb Faculty Scholar, University of California College of the Law, San Francisco; Senior Research Scholar, Center for Innovation (C4i) and AI Law & Innovation Institute (AILII). For helpful comments, questions, and conversations, I am grateful to Omri Ben-Shahar, Ryan Calo, John Crawford, Ben DePoorter, Scott Dodson, Robin Feldman, Lee Anne Fennell, Amanda Fish, Aziz Huq, Jared Mayer, Eric Martínez, Dave Owen, Austin Peters, Philip Petrov, Carla Reyes, and the law faculties at DePaul, Georgia State, Michigan State, Ohio State, Seton Hall, SMU, UC Law SF, University of Mississippi, and University of Washington.

¹ MARK ESSIG, EDISON & THE ELECTRIC CHAIR: A STORY OF LIGHT AND DEATH 215 (2003) (“Panic-stricken building owners took the law into their own hands and chopped the wires running over their rooftops. According to the *World*, some people grew so terrified that they threw out their telephones, ‘as if the little wires which connect them went straight to the river of death.’”).

own front doors, lest their deadbolts be electrified.² The gas lighting industry further stoked the public's fears, decrying the dangers of the "mystic nerve energy of electric wires."³ Judges described electricity as "a most subtle and dangerous agency" that "lurks unsuspected in the simple and harmless wire."⁴ And scholars worried that tort law was not up to the task of redressing its harms.⁵ But as the decades passed, electricity became ubiquitous, and the concern that tort law couldn't manage its harms vanished. Strict liability was roundly rejected, and negligence—the standard rule—ruled the day.⁶

A little over a century later, here I am, typing away on an electric computer, under the glow of electric lighting. My electric washing machine bangs away behind me, drowning out the music from my computer speakers. And because today my cold laundry room doubles as my home office, my feet rest on an electric floormat. I'm practically swimming in electricity—as are you, for that matter. Yet we hardly think about it. Negligence, it turns out, has worked just fine.

Today artificial intelligence is the new powerful-and-scary technology we're worried about. And our reactions to it—our excitement, and our fear—are all historically familiar. Also familiar is the worry that tort law might not be capable of redressing the harms caused by so new and mysterious a technology. Electricity was once seen as an invisible danger "lurk[ing] unsuspected" in wires.⁷ What dangers do we fear might lurk behind the screen when we engage with an AI chatbot? Soon, what anxieties might our robot housekeepers engender?

Today's AI fears, like the ones our ancestors held for electricity almost a hundred and fifty years ago,⁸ are not without basis. Several major companies

² See Sam Clarke & Di Cooke, *The 'Old AI': Lessons for AI Governance from Early Electricity Regulation*, EFFECTIVE ALTRUISM F. (Dec. 18, 2022), <https://forum.effectivealtruism.org/posts/k73qrimxcKtKZ4ng/the-old-ai-lessons-for-ai-governance-from-early-electricity-1> [perma.cc/ACY7-M27R] ("Mr. Edison has since declared that any metallic object—a doorknob, a railing, a gas fixture, the most common and necessary appliance of life—might at any moment become the medium of death." (internal citation omitted)).

³ Harold D. Wallace, Jr., *Electric Lighting Policy in the Federal Government, 1880-2016*, at 91 (July 19, 2018) (Ph.D. dissertation, University of Maryland) (ProQuest).

⁴ *Thornton v. Union Elec. Light & Power Co.*, 72 S.W.2d 161, 164 (Mo. Ct. App.1934).

⁵ Lester W. Feezer, *Tort Liability of Suppliers of Electricity*, 22 WASH. U. L. Q. 357, 359 (1937) ("The social needs of electricity in the normal life of everybody are increasing. This calls for more powerful currents, but at the same time for more careful safeguards. An accident happens, someone is injured or killed. Who shall bear the loss? What conduct will shift the loss to the electric company, and what factors are employed in determining that conduct? . . . [I]n electricity we are dealing with something new. There may be factors peculiar to the economic and social utility of the electrical industry. The 'mores' of the time and place, public opinion, the social sanctions, or whatsoever they may be called, may have a bearing upon allocating the risks which great utilities create by their existence.").

⁶ *Id.* at 358.

⁷ *Thornton*, 72 S.W.2d at 164.

⁸ The Electric Wire Panic of 1889 was set off by a series of electricity-related accidents, including the horrific electrocution of a lineman named John Feeks, who died suspended above a crowded street in a tangle of wires. See Peter Pappas, *The Dark Side of the Light Bulb*, THE FORGOTTEN FILES (Sep.

have conducted significant layoffs, citing AI-driven operational efficiency.⁹ AI deepfakes have created a whole new category of child sexual abuse material (CSAM).¹⁰ And in 2024 a fourteen-year-old boy died by suicide after a chatbot to which he was emotionally attached appeared to egg him on.¹¹ The list of potential harms is long.¹²

Some scholars are so concerned about the danger that artificial intelligence poses and believe that the risk of harm is so great—and tort law’s capacity for redress so lacking—that we need sweeping new liability rules.¹³ A large subset of these scholars believe categorical strict liability is the right rule.¹⁴ In their view the purveyors of artificial intelligence technologies ought, functionally, to be absolute insurers.

It’s easy to see what might motivate these proposals. Consider the latter two examples above: image-generation AI capable of creating CSAM, and chatbots that encourage suicidal ideations. Suppose the standard tort rules—negligence and products liability—fail to reach harms like these. What then? No wonder some call for heightened liability rules that promise to guarantee the kind of accountability that we fear the status quo won’t deliver.¹⁵

But just as that approach would have been a mistake a hundred years ago with electricity, it would be a mistake today with artificial intelligence. At the turn of the nineteenth century scholars feared negligence would fail to redress electrical injuries, thanks to hard-to-identify causes. But we managed, and workhorse doctrines like *res ipsa loquitur* earned their keep in the early days of electrical-injuries litigation.¹⁶ And courts turned out to be perfectly capable of

27, 2023), <https://forgottenfiles.substack.com/p/the-dark-side-of-the-light-bulb> [perma.cc/HT48-29NJ] (providing information surrounding the Electric Wire Panic of 1889).

⁹ Isobel O’Sullivan, *These Companies Have Already Replaced Workers with AI in 2025 and 2026*, TECH.CO, <https://tech.co/news/companies-replace-workers-with-ai> [perma.cc/8VD5-SXCL] (last visited Jan. 25, 2026).

¹⁰ *Artificial Intelligence (AI) and the Production of Child Sexual Abuse Imagery*, INTERNET WATCH FOUND., <https://www.iwf.org.uk/about-us/why-we-exist/our-research/how-ai-is-being-abused-to-create-child-sexual-abuse-imagery/> [perma.cc/X2RT-CSQ5].

¹¹ Kate Payne, *An AI Chatbot Pushed a Teen to Kill Himself, a Lawsuit Against Its Creator Alleges*, ASSOCIATED PRESS, <https://apnews.com/article/chatbot-ai-lawsuit-suicide-teen-artificial-intelligence-9d48adc572100822fdb3c90d1456bd0> [perma.cc/PCK6-4XGY] (Oct. 25, 2024).

¹² See *What Are the Risks of Artificial Intelligence?*, MIT AI RISK INITIATIVE, <https://airisk.mit.edu/> [perma.cc/MN8J-JF3G] (collecting over 1600 AI risks “categorized by their cause and risk domain”).

¹³ See *infra* note 148 and accompanying text (explaining that some scholars propose strict liability to combat new harms posed by AI).

¹⁴ See *infra* Part II.

¹⁵ David C. Vladeck, *Machines Without Principals: Liability Rules and Artificial Intelligence*, 89 WASH. L. REV. 117, 146 (2014) (contending that strict liability should be used in order to protect against AI harms).

¹⁶ See *Glasco Elec. Co. v. Union Elec. Light & Power Co.*, 61 S.W.2d 955, 957 (Mo. 1933) (“The *res ipsa loquitur* doctrine finds frequent application in cases against electric companies . . .”).

deciding what an electricity purveyor's duty of care ought to be.¹⁷ Negligence evolved.

This paper makes two primary arguments. First, we do not need new categorical liability rules to handle artificial intelligence—categorical strict liability least of all.¹⁸ Our default rules for unintentional torts—negligence and products liability—are high flexibility, high context doctrines that judges and juries will have no special difficulty applying to AI harms. Where these rules prove insufficient, private insurance, alongside narrow, tailored legislation, will be the right tools for the job.¹⁹ Second, even if special categorical rules just for AI harms were desirable, no modern theory of tort law justifies such a rule.²⁰

The article proceeds in four Parts. Part I explains why AI is unique among modern technologies; on what grounds some worry that tort law's status quo will struggle to handle its harms; and why AI activities and their resulting harms are not monolithic.²¹ Part II canvasses the various proposals others have made to deal with the supposed challenge of redressing AI harms; describes one of them in detail (categorical strict liability); and presents normative and doctrinal arguments against that proposal.²² Part III raises a more foundational objection to categorical strict liability for AI harms by making the case that no modern theory of tort law would support it.²³ Finally, Part IV briefly discusses insurance, targeted legislation, and the specter of catastrophic harms.²⁴

I. PROBLEMS, REAL AND IMAGINED

Let's start with the stakes. Artificial intelligence is already approaching ubiquity—and becomes more powerful and useful by the day. Today and tomorrow, what kinds of harms could AI actors²⁵ cause that tort law might plausibly be concerned with?

Count the kinds of harms a human can inflict, and you'll have your answer. As is true of humans, AI's capacity for harm is the photo negative of its capabilities. AI actors and their countless activities are highly differentiated—

¹⁷ *Geismann v. Missouri-Edison Elec. Co.*, 73 S.W. 654, 661 (Mo. 1903) (“[I]t was defendant’s duty, in the first place, to use every protection which was reasonably accessible to insulate its wires at the point of contact or injury in this case, and to use the utmost care to keep them so; and the fact of the death of Geismann is conclusive proof of the defect of the insulation, and negligence of the defendant . . .”).

¹⁸ See *infra* Part II.

¹⁹ See *infra* Part IV.

²⁰ See *infra* Part III.

²¹ See *infra* Part I.

²² See *infra* Part II.

²³ See *infra* Part III.

²⁴ See *infra* Part IV.

²⁵ This article uses the term “AI actor” to mean, simply, an AI entity that takes some action. Where it refers to some person or legal entity that deploys or uses an AI, it uses the term “deployer.”

the opposite of monolithic. The same is true for their potential harms.²⁶ Any coherent discussion of AI must account for the fact that “AI” could mean a hundred different things—and the phrase “AI harms” is nearly as sweeping and vague as “human harms.”

What’s more, some artificial intelligences can do things that humans can’t, further expanding AI’s range of potential harms. For example, no human voice actor is half as good as AI when it comes to mimicry, as Secretary of State Marco Rubio experienced in the first half of 2025.²⁷ Nor can any human animator create images as quickly or anonymously as AI can—a fact that has already led to the proliferation of a horrifying amount of AI-generated child pornography.²⁸ And where humans can do just one or two things at a time—a natural limit on our capacity for harm—some artificial intelligences are multi-tasking grandmasters.²⁹

In short, the stakes are high indeed. Even if artificial intelligence improves our lives and makes us safer overall, it appears likely that an increasing proportion of all harms will be caused by AI—and at least some such harms will be serious. Scholars and policymakers are right to consider carefully the liability questions artificial intelligence raises.

Of course, not all harms sound in tort. There are some AI harms that are indisputably beyond its capacity for redress—and so they are (mostly) beyond the scope of this article. Chief among these is the downward pressure on the demand for labor that AI might exert, and any subsequent adverse effects on employment.³⁰ Such risks may represent AI’s greatest threat to society, and so

²⁶ Work is underway to attempt to catalogue these harms. For example, MIT’s “AI Risk Repository” is a “living database” that includes over 1,600 “AI risks” to date. PETER SLATTERY ET AL., ARXIV, THE AI RISK REPOSITORY: A COMPREHENSIVE META-REVIEW, DATABASE, AND TAXONOMY OF RISKS FROM ARTIFICIAL INTELLIGENCE 2 (Mar. 2025), <https://doi.org/10.48550/arXiv.2408.12622>; MIT AI RISK INITIATIVE, *supra* note 12.

²⁷ See Humeyra Pamuk, *Rubio Impersonator Using AI Contacted Foreign Ministers, Cable Says*, REUTERS, <https://www.reuters.com/world/us/rubio-impersonator-used-ai-calls-foreign-ministers-cable-shows-2025-07-08/> [perma.cc/NLT2-Y286] (July 8, 2025).

²⁸ INTERNET WATCH FOUND., *supra* note 10.

²⁹ As of early 2025, ChatGPT had over 150 million weekly active users. See Akash Sriram, *Ghibli Effect: ChatGPT Usage Hits Record After Rollout of Viral Feature*, REUTERS, https://www.reuters.com/technology/artificial-intelligence/ghibli-effect-chatgpt-usage-hits-record-after-rollout-viral-feature-2025-04-01/?utm_source=chatgpt.com [perma.cc/4YFW-F4TW] (Apr. 1, 2025) (discussing the impact of ChatGPT’s image tool and demonstrating the average users weekly surpasses 150 million).

³⁰ See Daron Acemoglu, David Autor, Jonathon Hazell & Pascual Restrepo, *Artificial Intelligence and Jobs: Evidence from Online Vacancies*, 40 J. LAB. ECON. (SPECIAL ISSUE) S293, S296 (2022) (describing their findings as implying that “any positive productivity and complementarity effects from AI are at present small compared with its displacement consequences”). But see James Bessen, *Automation and Jobs: When Technology Boosts Employment* 34 (Mar. 2018) (unpublished manuscript), https://scholarship.law.bu.edu/faculty_scholarship/815/?utm_source=scholarship.law.bu.edu%2Ffaculty_scholarship%2F815&utm_medium=PDF&utm_campaign=PDFCoverPages [perma.cc/6NHG-5ZWN] (“Productivity-enhancing technology will increase industry employment if product demand is sufficiently elastic.”).

undoubtedly warrant serious analysis and discussion.³¹ But gravity alone cannot transport them into tort law's domain.³²

It's all the rest—the AI harms that do fall within tort law's analytical ambit—that this article and the views it engages are concerned with. As already noted, this category is as vast as AI is ubiquitous. Toast burnt by an AI toaster oven. Whiplash caused by a driverless car's overzealous braking. AI-generated revenge porn. Racial bias introduced by a poorly trained AI financial-credit analysis tool. From minor mistakes to systemic mayhem, AI will be capable of doing it all—just like we are.

A recent wave of scholarship has proposed that, in its current state, tort law is inadequate to meet the threat of AI's myriad potential harms.³³ Some support this view by arguing that AI in general is too dangerous to be handled by our standard rules.³⁴ Others take the view that AI has certain unique qualities that grant it an inappropriate degree of immunity to negligence and products liability.³⁵

Both grounds are mistaken. First, they fall prey to the essentialist fallacy, forgetting that artificial intelligences and their harms are not monolithic. And second, they presuppose doctrinal rigidity that neither negligence nor products liability possesses. These faulty premises arise from the same animating concern: that there is something qualitatively different about AI harms that demands special doctrinal attention. Section A of this Part begins by discussing whether that's right.³⁶ Section B and Section C then consider how negligence and products liability, respectively, might address AI harms.³⁷

A. Why Current-Generation AI Actors Are Unique

To understand what all the fuss is about, one needs to understand what sets modern artificial intelligence apart from earlier forms of AI. For decades we've had computer systems that could follow strict conditional rules—if X,

³¹ I am of the view that artificial intelligence will not harm the labor force in the medium or long term, but only time will tell.

³² Insofar as harms like these might be relevant to the discussion this article engages, it might be because some might propose heightened liability for AI harms as a means of stymieing AI innovation. Put another way, if, as a matter of policy, a government wanted to slow AI innovation in order to prevent a perceived harm to the labor force, then imposing liability-increasing rules like strict liability would likely further that agenda. That being said, contorting tort law to further some innovation-quashing policy aim would find no basis in any theory of tort law. Tort is simply the wrong tool for the job.

³³ See *infra* Part II for an in-depth discussion of the scholarship and proposals described in this paragraph.

³⁴ See *infra* Part II.

³⁵ See *infra* Part II.

³⁶ See *infra* Part I.A.

³⁷ See *infra* Part I.B–I.C.

then *Y*—to give the appearance of intelligence.³⁸ Some of these are highly complex and very useful. For example, expert systems have been used for decades to assist with everything from medical diagnosis to accounting to industrial design³⁹—and robot vacuums have been sweeping floors since 1996.⁴⁰

But what we think of today as artificial intelligence is very different. It is not, like its predecessors, a mere collection of if-then statements bound together by strings of code.⁴¹ Instead, modern artificial intelligence is built through advanced machine learning using neural networks.⁴²

Let's unpack these (and a few other) terms. Machine learning is a branch of computer science and artificial intelligence that uses statistical algorithms to help an AI model “learn”—which is to say, improve itself iteratively.⁴³ Machine learning enables software to become more effective without the need for a human to manually write or augment its code.⁴⁴ But it, too, has technically been around for decades.⁴⁵ What's different about modern artificial intelligence is the introduction of neural networks, and innovative ways of training them.⁴⁶

The neural network is what does the learning; functionally, it's the AI itself.⁴⁷ It comprises a huge number of interconnected nodes, which we can

³⁸ Vasudevan Swaminathan, *The Conundrum of Using Rule-Based vs. Machine Learning Systems*, ZUCISYS., <https://www.zucisystems.com/blog/the-conundrum-of-using-rule-based-vs-machine-learning-systems/> [perma.cc/59KM-3L6K] (“Rule-based systems are computer programs that use if-then rules to make decisions and perform tasks.”).

³⁹ Dorothy Leonard-Barton & John J. Sviokla, *Putting Expert Systems to Work*, HARV. BUS. REV., Mar.–Apr. 1988, at 91, <https://hbr.org/1988/03/putting-expert-systems-to-work> [perma.cc/9552-9W7P].

⁴⁰ Christopher White, *A Brief History of the Robot Vacuum Cleaner*, VACUUM WARS, <https://vacuumwars.com/history-of-the-robot-vacuum-cleaner/> [perma.cc/K53K-2HWW] (Apr. 11, 2025).

⁴¹ Swaminathan, *supra* note 38.

⁴² See IBM Data & AI Team, *AI vs. Machine Learning vs. Deep Learning vs. Neural Networks: What's the Difference?*, IBM, <https://www.ibm.com/think/topics/ai-vs-machine-learning-vs-deep-learning-vs-neural-networks> [perma.cc/EJ28-RMB4] (providing the differences between AI, machine learning, and neural networks and highlighting that neural networks “are a subset of machine learning and are the backbone of deep learning algorithms”).

⁴³ See *id.* (noting that “machine learning can leverage labeled data sets to inform its algorithm in supervised learning”); see also Zachary Henderson, *AI and Probabilistic Dispute Resolution*, 2025 WIS. L. REV. 215, 244 (describing the basic mechanisms involved in machine learning); Brian Shepard, *The Reasonableness Machine*, 62 B.C. L. REV. 2259, 2320 (2021) (“Machine learning is, at its basic core, a statistical approach that assesses progress based on how well the program increases predictive power . . . it ordinarily gets better as it gets more feedback data and further adjusts its statistical approach.”).

⁴⁴ See *What Is Machine Learning (ML)?*, GOOGLE CLOUD, <https://cloud.google.com/learn/what-is-machine-learning> [perma.cc/V5UZ-U5BU] (“Machine learning allows computer systems to continuously adjust and enhance themselves . . .”).

⁴⁵ See Alexander L. Fradkov, *Early History of Machine Learning*, 53 IFAC PAPERSONLINE 1385, 1385 (2020) (describing early machine learning systems dating as far back as the 1950s).

⁴⁶ See IBM Data & AI Team, *supra* note 42.

⁴⁷ See Henderson, *supra* note 43, at 245 (emphasizing that AI systems are based on “machine learning . . . using neural networks”).

think of as artificial neurons.⁴⁸ These nodes store numbers called “parameters”—*lots* of them.⁴⁹ The nodes of the largest AI models at the time of this writing are estimated to contain trillions of parameters.⁵⁰

Things get complicated from here. Because many scholars have written extensively on how deep learning and neural networks work, permit me to oversimplify things here to the extreme, even as I invite you to plumb my footnotes.⁵¹ In short, training a neural network involves iteratively exposing it to mountains of data,⁵² and using algorithms to update the neural network’s parameters over and over (and over) again.⁵³ If you were to look inside the neural network, you could actually see those billions or trillions of parameters.⁵⁴ But seeing them wouldn’t get you very far. To a human—indeed, even to an AI engineer—the parameters themselves convey no meaning. This, in essence, is the “black box” problem. No doubt fiddling with them would have *some* effect on how the AI works—but *what* effect would be anyone’s guess. Yet once the neural network is fully trained, the particular calibration of those numbers is what enables the model to work its magic.⁵⁵ Thanks to those near-countless parameters, the trained AI may be capable of responding and behaving in ways that seem downright human.⁵⁶

But human it is not. Kent Walker, President of Global Affairs at Google and Alphabet, rightly points out that modern AI is really just computational statistics.⁵⁷ This isn’t to diminish the technology or suggest that it’s any less

⁴⁸ IBM Data & AI Team, *supra* note 42.

⁴⁹ See Henderson, *supra* note 43, at 245–46 (explaining that parameters are the numbers stored by neural networks, that adjust based on data).

⁵⁰ See Aashiya Mittal, *How Many Parameters Does GPT-5 Have? Full Breakdown & Analysis*, ONGRAPH (Nov. 14, 2025), <https://www.ongraph.com/gpt5/> [perma.cc/3ZDL-4WZ3] (noting that OpenAI’s GPT-5 “is estimated to have between 2 trillion and 5 trillion parameters, based on current AI scaling trends and industry predictions”).

⁵¹ See, e.g., Ryan Calo, *Artificial Intelligence Policy: A Primer and Roadmap*, 51 U.C. DAVIS L. REV. 399, 404–07 (2017) (describing what AI is and how it has been deployed); Ryan McCarl, *The Limits of Law and AI*, 90 U. CIN. L. REV. 923, 924–37 (2022) (discussing what AI is, and how it has evolved over decades); Henderson, *supra* note 43, at 244–47 (presenting a primer on artificial intelligence, neural networks, and deep learning).

⁵² See IBM Data & AI Team, *supra* note 42.

⁵³ See Simeon Kostadinov, *Understanding Backpropagation Algorithm*, MEDIUM: TDS ARCHIVE (Aug. 8, 2019), <https://medium.com/towards-data-science/understanding-backpropagation-algorithm-7bb3aa2f95fd> [perma.cc/YZU8-43P4] (“In simple terms, after each forward pass through a network, backpropagation performs a backward pass while adjusting the model’s parameters (weights and biases).”).

⁵⁴ See Henderson, *supra* note 43, at 246 (“[M]odern AI models have billions, and sometimes even trillions, of parameters.”).

⁵⁵ See *id.* at 245–47.

⁵⁶ See *id.*

⁵⁷ Kent Walker, *Engaging in Frank Discussions About the Promise and Risks of Emerging Tech*, LINKEDIN (June 27, 2023), <https://www.linkedin.com/pulse/engaging-frank-discussions-promise-risks-emerging-tech-kent-walker/> [perma.cc/8YCK-K9AB].

powerful. But perspectives like Walker's are critical for distinguishing between what AI is—a powerful tool capable of producing remarkable probabilistic outputs—and what it is not: sentient.⁵⁸

This distinction is foundational to the debate this article joins. On one hand, artificial intelligence has become capable of taking actions that hitherto only humans could take. But on the other hand, when an artificial intelligence takes such actions, it does so mindlessly. Which is not to say it acts carelessly, lazily, recklessly, wantonly, or any other *-ly* we might attribute to a human's behavior. No, when an AI acts, it is mind-less, thought-less, emotion-less—despite every appearance to the contrary.⁵⁹

This raises the question: given all this, are AI actors also, in all cases, blame-less? If so, what then? This is one of the questions that undergirds this whole area of scholarship—and brings us to the doorstep of the question of whether tort law can redress harms caused by AI actors. The next two subparts consider how two obvious legal theories—negligence and products liability—might handle AI harms.⁶⁰

B. Negligence Can Handle AI Harms

Many argue that negligence is insufficient to redress harms caused by AI actors—either because AI actors cannot themselves be negligent, or because proving negligence against a manufacturer, seller, or user would be too difficult.⁶¹ I find this argument unavailing. It's true AI actors cannot themselves be negligent—more on this point in a moment—but the legal entities that deploy them certainly can be. Negligence is a high flexibility, high context doctrine grounded in societal views on what kinds of behavior are, and are not, reasonable.⁶² There's no reason to think the unreasonable deployment of a harm-causing AI actor would be insulated from a breach-of-duty assertion under a

⁵⁸ Not yet at least.

⁵⁹ This may change in the future, but it is indisputably true for now. *If* it changes, the important first questions will not be of law (let alone tort), but of philosophy and ethics.

⁶⁰ See *infra* notes 61–148 and accompanying text.

⁶¹ See, e.g., Anat Lior, *AI Entities as AI Agents: Artificial Intelligence Liability and the AI Respondeat Superior Analogy*, 46 MITCHELL HAMLINE L. REV. 1043, 1057 (2020) (“The lack of foreseeability, AI entities’ varying degrees of autonomy, and the absence of complete human control with regards to the potential behavior of AI entities leads to difficulty in establishing a legal nexus of causation between the victim and the tortfeasor . . .”).

⁶² The “reasonable person” standard is foundational to the doctrine of negligence. See *Balt. & P.R. Co. v. Jones*, 95 U.S. 439, 441–42 (1877) (“Negligence is the failure to do what a reasonable and prudent person would ordinarily have done under the circumstances of the situation, or doing what such a person under the existing circumstances would not have done.”); RESTATEMENT (THIRD) OF TORTS: LIAB. FOR PHYSICAL AND EMOTIONAL HARM § 3 (A.L.I. 2010) (“A person acts negligently if the person does not exercise reasonable care under all the circumstances.”).

negligence analysis.⁶³ Accordingly, where negligence fails to redress a given AI harm, it will be because there was no negligence to speak of.

1. AI Actors Are Not Susceptible to Negligence Claims . . .

Let's quickly dispose of the obvious: as a matter of law, AI actors can't be negligent.⁶⁴ This is just as true for modern, neural-network-based artificial intelligence as it is for older, simpler kinds of AI. For one thing, they have no legal status. For another, the linchpin of negligence analysis is the reasonable person standard. From the third restatement: "A person acts negligently if the person does not exercise reasonable care under all the circumstances."⁶⁵ The emphasis on "person" says it all: negligence law is about assessing human choices.⁶⁶ True, corporate persons can be held negligent too, both vicariously⁶⁷ and directly.⁶⁸ But in all cases, a finding of negligence is rooted in a failure of human decision-making.⁶⁹

But forget what the law is today. Why not shoehorn AI actors into the negligence framework? After all, this is one of the things that some of the scholarship advocating AI personhood appears to seek.⁷⁰

Try all you want; the shoe just doesn't fit. While current AI systems can *project* human-like rationality, they possess no such thing.⁷¹ An AI system's output, however impressive, is the product of computational statistics and probability—not reasoned thought.⁷² Given this, what would it mean to assign an AI system a duty of care? And even if we managed to get our heads around

⁶³ The four elements of negligence, stated plainly, are duty, breach, causation, and damages. *See* WILLIAM M. KUNSTLER, *LAW OF ACCIDENTS* 9 (1954) (describing the four elements of negligence); *see also* RESTATEMENT (SECOND) OF TORTS § 281 (A.L.I. 1965) (describing the elements of negligence).

⁶⁴ RESTATEMENT (THIRD) OF TORTS: LIAB. FOR PHYSICAL AND EMOTIONAL HARM § 3 (looking at the conduct of a "person" to define negligence).

⁶⁵ *Id.* (endorsing the "reasonable person" standard for negligence liability).

⁶⁶ David G. Owen, *The Five Elements of Negligence*, 35 HOFSTRA L. REV. 1671, 1675 (2007).

⁶⁷ *See* Phila. & Reading R.R. Co. v. Derby, 55 U.S. (14 How.) 468, 487 (1852) (affirming judgment finding railroad liable for the negligence of one of its conductors).

⁶⁸ *See* Norfolk & W. Ry. Co. v. Ayers, 538 U.S. 135, 144, 160 (2003) (affirming finding of negligence against railroad for asbestos-related injuries).

⁶⁹ Owen, *supra* note 66, at 1675 ("At bottom, negligence law assesses human choices to engage in harmful conduct as proper or improper.").

⁷⁰ *See, e.g.,* Brian Gannon, *Autonomous Machines and the Law*, 2 CITY L. REV. 161, 169 (2020) ("Conferring legal personality on AI machines would resolve the agency question since as a principal the machine assumes rights and obligations, just as a legal corporation does. This proposal extends well beyond the boundaries of negligence.").

⁷¹ *See* Ira Winkler, *Artificial Intelligence Is Just Math: Looking Beyond the Hype*, CYE (Mar. 24, 2025), <https://cyesec.com/blog/artificial-intelligence-is-just-math-looking-beyond-hype> [perma.cc/Q3WV-GGL3] ("Essentially, AI is a special set of mathematical algorithms that is more robust than traditional algorithms.").

⁷² *Id.*; *see also* Walker, *supra* note 57.

duty and breach, what then? AI systems have no legal status; they can't be sued.⁷³ Without some way of tracing responsibility for the harm back to some legal entity, accusing an AI actor of negligence would bear as little fruit as accusing a chainsaw. Of course, the idea of a negligent chainsaw wouldn't even cross our minds; they make no decisions, let alone act autonomously.⁷⁴ If we were injured by one, we would sue some person or company for negligence, or we would sue its manufacturer or seller, on products-liability grounds.⁷⁵

The upshot is that if we want to raise a negligence claim, it will need to be against some legal entity associated with the AI actor that caused the harm. Let's discuss that next.

2. . . . but the Legal Entities That Deploy Them *Can* Be Negligent

Even though AI actors cannot be held negligent, their legal-entity deployers certainly can be. For critics, this may be small comfort. After all, just because the makers or users of AI *can* be negligent doesn't mean they will be found negligent when, for example, a self-driving car injures a pedestrian.

Perhaps. But negligence is a high flexibility, high context doctrine—and courts have ample tools to assign negligence correctly.⁷⁶ In the case of AI harms, the analytical focus will simply shift backward in time a bit. Rather than asking if the AI actor behaved negligently (a legal and technological impossibility we've already dismissed), courts will ask whether deploying the AI in the first place breached the required standard of care.

Let's pause again to define some terms I've already used several times. I'm using "AI actor" to mean an AI capable of action, and "deployer" as broad shorthand for any legal entity that takes some step toward putting an AI actor out into the world. If a manufacturer places an AI actor in the stream of commerce, that would be a deployment. Same thing when a third-party vendor sells one to an end user—or when that end user actually uses the AI actor to complete a task. (Interpret "use" broadly, too. The deployment may amount to little more than activating the AI, or setting it loose, so-to-speak. AI agents are increasingly capable of sustained, self-directed activity.)⁷⁷

Any such deployment is analyzable through the lens of negligence. Take the sale of an AI-powered vehicle marketed as capable of unsupervised, fully

⁷³ Self-aware and rights-deserving artificial intelligences are probably in humanity's future. But by all accounts, that's still decades or centuries away.

⁷⁴ Though if one did, our products-liability victory would be assured!

⁷⁵ See Vladeck, *supra* note 15, at 132 (noting that "cases involving complex products are now typically viewed through the lens of products liability law").

⁷⁶ RESTATEMENT (THIRD) OF TORTS: LIAB. FOR PHYSICAL AND EMOTIONAL HARM § 3 (A.L.I. 2010).

⁷⁷ See *What Is an AI Agent?*, GOOGLE CLOUD, <https://cloud.google.com/discover/what-are-ai-agents> [perma.cc/BHR5-98P5] (Dec. 4, 2025).

autonomous driving. Suppose that despite its marketing, it is in fact *not* fully autonomous. It can't see well, routinely stops steering or braking, and the end-user frequently has to take the wheel without notice. If such a car causes harm on account of these failings, then the act of selling the car is almost certainly breach-of-duty enough to support a negligence suit. (Let's put aside products liability for now—we'll discuss it separately and at length later.)⁷⁸

But suppose the car *is* fully autonomous, just accident-prone: it has an accident rate that is roughly double the rate of accidents caused by drivers aged sixteen to seventeen—by far the statistically worst human drivers.⁷⁹ Suppose *this* autonomous vehicle causes an accident, and the victim sues the manufacturer. What then? Is deployment in this case also enough to show breach? Very likely yes. To escape negligence, the seller would need to show that deploying the car—selling it into the market—was reasonable.⁸⁰ That will be tough, given that this car isn't nearly as safe as even the worst category of human drivers. A fact finder is likely to conclude that selling a car that comes packaged with so unsafe an AI driver is a breach of the duty of care.⁸¹ After all, it wouldn't be difficult to argue that by selling this car, the manufacturer made the roads less safe—and if all cars were replaced with this one, the accident rate across the country would skyrocket.

Interestingly, neither of the above scenarios represent the reality of unsupervised, autonomous vehicles in the United States. Quite the opposite: the unsupervised, autonomous cars on the road today have accident rates that are *far* below the accident rates caused by human drivers.⁸² When we take this scenario as our hypothetical, the negligence analysis gets quite a bit trickier. Does selling a car like this violate the duty of care?

Maybe yes, maybe no. Even supposing factfinders are willing to conclude “no breach” when an AI driver is “safer” than a human driver, defining what “safer” means is surprisingly hard. Let's get concrete. Waymo's driverless taxis are currently carting thousands of people around San Francisco, Los Angeles,

⁷⁸ See *infra* Part I.C.

⁷⁹ See Brian C. Tefft, *Rates of Motor Vehicle Crashes, Injuries and Deaths in Relation to Driver Age, United States, 2014-2015*, AAA FOUND. FOR TRAFFIC SAFETY (June 2017), <https://aaaafoundation.org/rates-motor-vehicle-crashes-injuries-deaths-relation-driver-age-united-states-2014-2015/> [perma.cc/6DDY-RE9Z].

⁸⁰ RESTATEMENT (THIRD) OF TORTS: LIAB. FOR PHYSICAL AND EMOTIONAL HARM § 3.

⁸¹ See RESTATEMENT (SECOND) OF TORTS § 281 (A.L.I. 1965) (providing the elements of negligence).

⁸² See *Research & Discoveries (R&D): Waymo Reduces Crash Rates Compared to Human Drivers Over 7+ Million AV Miles*, AUTONOMOUS VEHICLE INDUS. ASS'N (Jan. 17, 2024), <https://theavindustry.org/resources/blog/waymo-reduces-crash-rates-compared-to-human-drivers> [perma.cc/SS5A-Q87L] (highlighting that Waymo AI drivers demonstrate an 85% reduction in injury-causing crashes, and a 57% reduction in all police-reported crashes).

Phoenix, Atlanta, and Austin,⁸³ and they're doing so at accident rates far below those associated with human drivers⁸⁴ (I personally always prefer Waymo to a human-driven taxi or rideshare). Let's say that Waymo AI drivers are only ten percent as likely as a human driver to *cause* an accident (this closely resembles the actual data).⁸⁵ Given these impressive statistics, one might argue that the act of deploying a Waymo AI vehicle cannot be negligent with respect to any accidents it gets into. After all, how could making the roads *safer* breach the seller's duty of care?

But let's dig a bit deeper. Is the Waymo AI driver safer than *reasonable* human drivers, or is it only safer than negligent ones? The answer isn't obvious; we'd need to parse out fault and no-fault accidents in our human-driver accident statistics.

Bearing this wrinkle in mind, what if it turns out that *all* car accidents involving only human drivers are caused by some driver acting negligently? In other words, suppose human negligence was a factor for 100% of human-only car accidents. In that scenario, Waymo would be a bit less impressive. True, we know that Waymo AI drivers cause ninety percent fewer accidents than human drivers.⁸⁶ But statistically, Waymo AI drivers would drive worse than reasonable human drivers: Waymo would still cause *some* car accidents, whereas in this (implausible) hypothetical, reasonable human drivers cause none. This would invite the conclusion that the Waymo AI driver does not drive as safely as a reasonable driver would—and so its deployment might arguably breach the manufacturer's duty of care.⁸⁷

On the other hand, suppose that only seventy percent of accidents are caused by negligent human drivers; the remaining thirty percent are no-fault, caused by human drivers driving reasonably. In this case, even if we removed the negligence-caused accidents from the comparison, the Waymo AI driver would still be statistically safer: it causes a third the number of accidents that reasonable human drivers cause. In this scenario, Waymo is especially impressive. Not only is the Waymo AI driver safer than negligent human drivers; it's arguably safer than two thirds of reasonable human drivers as well.

Whatever the empirical answers to these questions, for our purposes it is enough to show that it's possible to analyze deployments of AI for negligence. This same analytical framework can be applied to situations where end users deploy AI actors, too. For example, someone who owns a sophisticated AI dog nanny might reasonably deploy it to feed, water, and pet his dogs while he's at

⁸³ *Where Waymo Is Driving*, WAYMO, <https://waymo.com/rides/#rides-map> [perma.cc/ZV2X-U7ER].

⁸⁴ AUTONOMOUS VEHICLE INDUS. ASS'N, *supra* note 82.

⁸⁵ *Id.*

⁸⁶ *Id.*

⁸⁷ See KUNSTLER, *supra* note 63, at 9 (providing the elements of negligence).

work. But he probably breaches his duty of care if he deploys it to babysit his toddler while he goes to the ballpark.⁸⁸

Before we turn to products liability, there is one other story to tell about AI and negligence. Though negligence cannot directly reach AI actors, the existence of AI actors is certain to have a profound effect on what human actions constitute negligence.

3. The Indirect Effects of AI on Negligence and the Reasonable Person Standard

Recall that driverless taxis are already a bustling reality, and the statistical evidence suggests they are operating far more safely than human drivers.⁸⁹ In fact recent reports suggest that pedestrians, cyclists, and pets stand particularly to gain: Waymo driverless taxis can see them and react to them far better than their human-driver counterparts.⁹⁰ This is great news for safety—and has interesting potential implications for the reasonable-person standard.

Ryan Abbott presented the bull case back in 2018.⁹¹ He argued that the advent of AI actors with capabilities superior to humans will *directly* alter the reasonable-person standard by raising the bar for reasonableness to match the AI actor's capabilities.⁹² So, for example, if a driverless car is capable of superior accident avoidance, then that capability will establish a new yardstick against which the reasonableness of a human driver's actions will be measured.⁹³

I am less convinced than Abbott that the reasonable person standard will or should be *directly* altered to align with the superior capabilities of AI actors.⁹⁴ But I have no doubt AI will change the standard *indirectly*. For example, it seems likely that, one day soon, choosing to manually operate a vehicle de-

⁸⁸ See RESTATEMENT (THIRD) OF TORTS: LIAB. FOR PHYSICAL AND EMOTIONAL HARM § 3 (A.L.I. 2010) (providing the standard of care under negligence).

⁸⁹ AUTONOMOUS VEHICLE INDUS. ASS'N, *supra* note 82.

⁹⁰ See Kristofer D. Kusano et al., *Comparison of Waymo Rider-Only Crash Rates by Crash Type to Human Benchmarks at 56.7 Million Miles*, 26 TRAFFIC INJ. PREVENTION S8, S8 (2025) <https://doi.org/10.1080/15389588.2025.2499887> (“Of the crash types, V2V Intersection crash events represented the largest total crash reduction, with a 96% reduction in Any-injury-reported (87–99% confidence interval) and a 91% reduction in Airbag Deployment (76–98% confidence interval) events. Cyclist, Motorcycle, Pedestrian, Secondary Crash, and Single Vehicle crashes were also statistically reduced for the Any-Injury-Reported outcome.” (emphasis added)).

⁹¹ Ryan Abbott, *The Reasonable Computer: Disrupting the Paradigm of Tort Liability*, 86 GEO. WASH. L. REV. 1, 1 (2018).

⁹² *Id.* at 37 (“In practice, this would mean that instead of judging a defendant’s action against what a reasonable person would have done, the defendant would be judged against what a computer would have done.”).

⁹³ *Id.* (“[T]oday a defendant might not be liable for striking a child running in front of their car if a reasonable driver would not have been able to stop immediately. But that person would soon be liable under the exact same circumstances if an automated car would have prevented the injury.”).

⁹⁴ Ask me again once we’ve entered the era of AI personhood.

spite the push-button availability of AI assistance will be deemed unreasonable for purposes of establishing negligence.

The hypo is easy to construct; let's expand on a scenario that Abbott discussed.⁹⁵ Suppose you buy a new car with a built-in AI driver. Reliable studies cited in your new car's marketing materials show that the AI driver can react to the sudden appearance of a pedestrian in mere milliseconds. Further suppose that you're driving the car with the AI driver turned off—and you unintentionally strike and kill a child chasing a basketball into the street. Let's say you weren't speeding, and your reaction time was exemplary: you slammed on the brakes faster than 90% of human drivers could have. Yet if the AI system had been engaged, its reaction time would have been so superior to yours that the child would not have been struck at all. Arguably, not using the AI driver—failing to *deploy* it, to reprise our earlier term—in a place where it was foreseeable that children might dash into the street, was unreasonable, and a breach of your duty of care.⁹⁶

Play this out across any number of situations in which an AI actor becomes safety-superior to humans. It's easy to imagine how the existence of AI actors with documented track records superior to humans will alter the reasonable person standard, and in turn, alter acceptable human behavior.

The upshot of this whole negligence discussion is twofold. First, until AI actors gain something like legal personhood, holding them negligent is probably a non-starter. Second, we *will* still be able to hold deployers accountable for negligence in the appropriate case. As AI actors improve and eventually surpass human capacities for safety, showing breach of duty on the deployer's part will become increasingly difficult—but rightly so.

C. AI Has No Products Liability Immunity

There's certainly nothing improper about trying to apply products liability to harms caused by AI products.⁹⁷ But some worry—incorrectly in my view—

⁹⁵ Abbott, *supra* note 92, at 22.

⁹⁶ This hypothetical does not imply that it will soon be negligent to drive a car without AI capabilities. State and federal legislation will continue to define what qualities and features vehicles must have in order to be deemed road-worthy. Instead, the linchpin of this hypothetical is that, at the moment of causation, the safety-maximizing decision—activating the AI driver instead of manually driving—was nearly costless.

⁹⁷ Even in the case of pure software applications of AI, courts appear to be increasingly willing to entertain products liability suits. *See, e.g.,* *Garcia v. Character Techs., Inc.*, 785 F. Supp. 3d 1157, 1180 (M.D. Fla. 2025), *motion to certify appeal denied*, No. 6:24-CV-1903, 2025 WL 2581834 (M.D. Fla. July 15, 2025) (“Even though Sewell may have been ultimately harmed by interactions with Character A.I. Characters, these harmful interactions were only possible because of the alleged design defects in the Character A.I. app. Accordingly, Character A.I. is a product for the purposes of Plaintiff's product liability claims so far as Plaintiff's claims arise from defects in the Character A.I. app rather than ideas or expressions within the app.”).

that when it comes to neural-network-based AI, doing so will be impracticable. In the United States, there are three ways a product can be deemed defective under products liability.⁹⁸ First, a product contains a *manufacturing* defect “when the product departs from its intended design.”⁹⁹ Second, a product contains a *design* defect when some reasonable alternative design could have prevented or reduced the likelihood of the product’s foreseeable risks, and omitting that design rendered the product unreasonably unsafe.¹⁰⁰ Finally, a product can be defective on account of *inadequate instructions or warnings*.¹⁰¹

To see why some fear it will be too difficult to apply products liability to AI harms, let’s consider these three potential defects in turn, starting with design defects. But before we do, a critical qualification: at most, these concerns apply to modern artificial intelligence that is susceptible to the “black box” problem described in the last subpart.¹⁰² Artificial intelligence based on *traditional* programming—the kind that’s been around for decades—is no more immune to a products liability claim than is a lawnmower.¹⁰³ Accordingly, when this subpart refers to AI actors or AI harms, keep in mind we’re focused on a narrow class of AI: modern, probabilistic artificial intelligence, and not old-school, deterministic AI.

⁹⁸ RESTATEMENT (THIRD) OF TORTS: PRODS. LIAB. § 2(a)–(c) (A.L.I. 1998).

⁹⁹ *Id.* § 2(a). Manufacturing defects are typically subject to strict liability. *See id.* (noting that the deviation is a manufacturing defect “even though all possible care was exercised in the preparation and marketing of the product”).

¹⁰⁰ *Id.* § 2(b) (explaining a design is defective “when the foreseeable risks of harm posed by the product could have been reduced or avoided by the adoption of a reasonable alternative design . . . and the omission of the alternative design renders the product not reasonably safe”). Or, in some jurisdictions, when the product is more dangerous than an ordinary consumer would have expected it to be. *See Ray ex rel. Holman v. BIC Corp.*, 925 S.W.2d 527, 529 (Tenn. 1996) (noting the “consumer expectation” test is “defined generally as, whether the product’s condition poses a danger beyond that expected by an ordinary consumer with reasonable knowledge, [and] has been employed by many states”); *see also* RESTATEMENT (SECOND) OF TORTS § 402A (A.L.I. 1965).

¹⁰¹ RESTATEMENT (THIRD) OF TORTS: PRODS. LIAB. § 2(c) (explaining a product might be defective “because of inadequate instructions or warnings when the foreseeable risks of harm posed by the product could have been reduced or avoided by the provision of reasonable instructions or warnings . . . and the omission of the instructions or warnings renders the product not reasonably safe”).

¹⁰² *See supra* Part I.B.

¹⁰³ Except on grounds it is not a product at all, as has been the case for much software. *See Winter v. G.P. Putnam’s Sons*, 938 F.2d 1033, 1036 (9th Cir. 1991) (declining to treat the ideas and expressions in a book as a product for products-liability purposes and suggesting computer software might warrant similar treatment). Notably, not only are there indications that courts are changing course here—*see Garcia v. Character Techs., Inc.*, 785 F. Supp. 3d 1157, 1180 (M.D. Fla. 2025), *motion to certify appeal denied*, No. 6:24-CV-1903, 2025 WL 2581834 (M.D. Fla. July 15, 2025)—but there is an ongoing bipartisan initiative to ensure that artificial intelligence is treated as a product for products liability purposes. *See Dick Durbin & Josh Hawley, Aligning Initiatives for Leadership, Excellence, and Advancement in Development (AI LEAD) Act: One Pager*, S. COMM. ON THE JUDICIARY, <https://www.judiciary.senate.gov/imo/media/doc/One-Pager%20-%20AI%20LEAD%20Act.pdf> [perma.cc/L5VS-3NJU] (“The *AI LEAD Act* would classify AI systems as products and create a federal cause of action for products liability claims to be brought when an AI system causes harm.”).

The argument that products liability will struggle to redress the harms caused by neural-network-based AI is strongest with respect to design and manufacturing defects. Recall how neural networks are developed.¹⁰⁴ The AI undergoes extensive algorithmic training by ingesting vast amounts of data.¹⁰⁵ During this iterative process, the algorithm updates the neural network's parameters—millions, billions, or even trillions of them—countless times.¹⁰⁶ Once training is complete, the AI will be capable of responding to inputs in remarkable ways—say, by being a friendly conversation partner, or by piloting a vehicle.

As we discussed earlier, if you were to look under the hood, you would see little more than a giant array of seemingly random numbers—those trillion parameters we mentioned in Part I.A.¹⁰⁷ Recall that even for the AI's software engineers, that vast collection of numbers is all but meaningless. There is no way to discern the system's likely functionality from even a careful study of them. This is why we call neural networks “black boxes.”¹⁰⁸ We'll return to this in a moment.

Now suppose an AI causes harm. Perhaps an AI chatbot convinces a human with a shellfish allergy to sample chapulines—a kind of grasshopper—on a trip to Central America, sending him into anaphylactic shock.¹⁰⁹ Or perhaps a self-driving car slams on its brakes and skids into a parked car after a large but harmless party balloon drifts into the road. Based on these actions, how might one argue that these AI actors were defectively designed?

To begin with, one would need to show that the risk of those harms was foreseeable.¹¹⁰ In each of these cases, foreseeability is at least plausible.¹¹¹

¹⁰⁴ See *supra* Part I.A.

¹⁰⁵ See IBM Data & AI Team, *supra* note 42.

¹⁰⁶ See Kostadinov, *supra* note 53.

¹⁰⁷ See *supra* Part I.A.

¹⁰⁸ See Vikas Hassija et al., *Interpreting Black-Box Models: A Review on Explainable Artificial Intelligence*, 16 COGNITIVE COMPUTATION 45, 47 (2023) (“A black-box model in XAI refers to a machine learning model that operates as an opaque system where the internal workings of the model are not easily accessible or interpretable. These models make predictions based on input data, but the decision-making process and reasoning behind the predictions are not transparent to the user.”).

¹⁰⁹ Chapulines are a kind of grasshopper eaten in certain parts of Central America. William N. Sokol, Sabina Wünschmann & Sayeh Agah, *Grasshopper Anaphylaxis in Patients Allergic to Dust Mite, Cockroach, and Crustaceans: Is Tropomyosin the Cause?*, 119 ANNALS ALLERGY, ASTHMA & IMMUNOLOGY 91, 91 (2017). It is common for people with crustacean allergies to also be allergic to edible insects like grasshoppers, on account of their having similar biochemistry. *Id.* at 91–92.

¹¹⁰ See RESTATEMENT (THIRD) OF TORTS: LIAB. FOR PHYSICAL AND EMOTIONAL HARM § 3 (A.L.I. 2010) (noting factors to consider in determining whether conduct lacks reasonable care includes the foreseeable likelihood that harm would result from the conduct). Note that, for brevity, I'm intentionally leaving out analysis of the minority-rule consumer-expectations test described in the Second Restatement.

¹¹¹ Some scholars believe foreseeability will be a consistent sticking point for AI actors. See, e.g., Lior, *supra* note 61, at 1057 (highlighting that AI decisions lack foreseeability, creating difficulty with regard to establishing causation between tortfeasors and victims, as well as between damages and

With respect to the chatbot, it's certainly foreseeable that a user with a food allergy might ask it for food recommendations.¹¹² After all, the manufacturers of these chatbots have explicitly designed them to be able to respond to virtually any input¹¹³ (whether it's foreseeable that a person might *rely* on the chatbot to proactively warn about potential food allergens is a closer call that implicates the chatbot's instructions and warnings—more on that later). As for the driverless car and the balloon, it's certainly foreseeable that large objects might suddenly enter the car's path—and foreseeable that slamming on the brakes to avoid them could result in skidding and property damage.

Some worry that foreseeability will be difficult or impossible in many cases.¹¹⁴ I'm unconvinced. The AI manufacturers know that the technology is capable of an all but unlimited range of outputs and know that the technology will hallucinate or produce unexpected outputs at least some of the time. Accordingly, judges and juries alike will, in my view, set the foreseeability bar extremely low where AI harms are involved. Of course, foreseeability is only one piece of the design-defect analysis. At least two Olympic high jumps remain: proof that a reasonable alternative design was possible, and proof that failure to use that alternative design rendered the AI actor unreasonably unsafe.¹¹⁵

These hurdles bring us back to the “black box” problem. How would one go about showing that there was a reasonable alternative design for the AI actor's neural network? To begin with, there is no AI expert who could credibly testify that tweaking some of the AI's parameters in a particular way would have made it safer, let alone prevented the harm at issue.¹¹⁶ Perhaps one could argue that the AI needed more training. But how much additional training would have been enough to make the AI acceptably safe—and what kind? Models can *always* be trained more for a price—and the cost-benefit tradeoffs

liable party). I think this is a bit overstated. As robust and occasionally creative as AI actors are, their behaviors nonetheless tend to follow some discernable (and therefore predictable) logic.

¹¹² Because chatbots generate language, there are interesting First Amendment questions, too. See the *Encyclopedia of Mushrooms* case, *Winter v. G.P. Putnam's Sons*, 938 F.2d 1033, 1037 (9th Cir. 1991) (“[T]here is nothing inherent in the role of publisher or the surrounding legal doctrines to suggest that . . . a duty [to investigate the accuracy of the contents of the books they publish] should be imposed on publishers.”). See generally Peter N. Salib, *AI Outputs Are Not Protected Speech*, 102 WASH. U. L. REV. 83 (2024) (arguing that AI outputs are not protected speech); Eugene Volokh, Mark A. Lemley & Peter Henderson, *Freedom of Speech and AI Output*, 3 J. FREE SPEECH L. 651 (2023) (arguing the opposite).

¹¹³ See, e.g., Mortuza Hossain, *How to Use Grok AI: A Comprehensive Guide*, DORIK (Feb. 3, 2025), <https://dorik.com/blog/how-to-use-grok-ai> [perma.cc/W2W4-Y7NM] (explaining how to use Grok AI, and noting that, at the time of writing, the textbox of the Grok interface invites the user to “Ask anything”).

¹¹⁴ Lior, *supra* note 61, at 1057–58.

¹¹⁵ RESTATEMENT (THIRD) OF TORTS: PRODS. LIAB. § 2(b) (A.L.I. 1998).

¹¹⁶ See Hassija et al., *supra* note 108, at 46 (“AI algorithms suffer from opacity i.e., the situation in which a system is unable to offer any reason or suitable explanation involved behind its decisions . . .”).

inherent in alternative designs are always relevant to the design-defect analysis.¹¹⁷ None of this is to say that a reasonable alternative design for an AI system could never be proposed.¹¹⁸ Indeed, odds are we'll soon see expert witnesses battling out whether more training would or would not have prevented this or that harm. But current AI systems do have at least a small measure of natural immunity to reasonable-alternative-design arguments.¹¹⁹

Now suppose we've managed to overcome both foreseeability and the reasonable alternative design requirement. We would still need to show that not adopting the alternative design made the AI actor unreasonably unsafe.¹²⁰ This is likely to be the highest hurdle of all, for reasons that will be familiar from our negligence discussion.

Once again, consider driverless cars. As we discussed earlier, driverless cars are already demonstrably safer than human drivers.¹²¹ So even if driverless cars *do* occasionally cause harm that a human driver might not have caused, this prong of our design-defect analysis will be very difficult to win on.¹²² Put another way, if AI driver 1.0 is already safer than human drivers, how might one show that not adopting improved AI driver 1.1 renders the previous version unreasonably unsafe?¹²³ Either version is safety-enhancing compared to the human driver it displaces. Of course, this example may be less indicative of the challenge of proving design defects in the case of AI harms, and more a reminder that not all harms are tortious.¹²⁴

¹¹⁷ RESTATEMENT (THIRD) OF TORTS: PRODS. LIAB. § 2 cmt. a (A.L.I. 1998) (“Products are not generically defective merely because they are dangerous. Many product-related accident costs can be eliminated only by excessively sacrificing product features that make products useful and desirable. Thus, the various trade-offs need to be considered in determining whether accident costs are more fairly and efficiently borne by accident victims, on the one hand, or, on the other hand, by consumers generally through the mechanism of higher product prices attributable to liability costs imposed by courts on product sellers.”).

¹¹⁸ One plausible argument might go like this: Newly-released version 1.2 of a given AI actor represents a reasonable alternative design, and version 1.1—the AI actor at issue—is defective by comparison. But even if this argument were successful, it might end up being a one-trick pony. Depending on the circumstances and preclusion rules of the jurisdiction, a ruling based on this kind of argument might make future arguments that version 1.2 is defectively designed more difficult.

¹¹⁹ See Brandon W. Jackson, *Artificial Intelligence and the Fog of Innovation: A Deep-Dive on Governance and the Liability of Autonomous Systems*, 35 SANTA CLARA HIGH TECH. L.J. 35, 59 (2019) (“Traditional concepts of product liability will likely fail to provide relief because a manufacturing defect cannot be identified, and the reasonableness of a jury to infer the cause of the harm will become increasingly attenuated.”).

¹²⁰ See RESTATEMENT (THIRD) OF TORTS: PRODS. LIAB. § 2(b) (explaining that a defective design is one that could have been avoided via a reasonable alternative design).

¹²¹ See *supra* Part I.B.

¹²² Driverless cars that cause systematic and predictable harms are probably another matter—but that situation is probably a better fit for negligence; we discussed precisely this in Part I.B.

¹²³ And then there's the repair rule to worry about. See FED. R. EVID. 407 (limiting the admission of evidence about remedial measures to prove, among other things, “a defect in a product or its design”).

¹²⁴ This article elaborates on this point in Part III.

The chatbot hypothetical is more nuanced. Whereas our driverless-car defendant has objective evidence of safety to point to—safety data for driverless vs. human-driven cars—no such comparison exists for AI vs. human conversationalists. In a vacuum, the chatbot is much more susceptible to design-defect claims, on grounds failure to improve the chatbot might render it unreasonably unsafe.

But chatbots do not exist in a vacuum and are a kind of AI actor for which instructions and warnings likely will (and in my view, should) go a long way. Virtually all foreseeable chatbot harms result from improper human reliance on what the chatbot says.¹²⁵ So clear reminders that AI chatbots hallucinate,¹²⁶ can be mistaken,¹²⁷ occasionally say disturbing things,¹²⁸ and should never be relied on without verification, might greatly reduce their risk of harm.¹²⁹ The existence of robust warnings and reminders will make proving that the chatbot is unreasonably unsafe in the absence of the alternative design very hard.¹³⁰

As difficult as it will be to show that an AI actor was defectively designed, it will be even harder to show one was defectively manufactured.¹³¹ To show a manufacturing defect would require (1) identifying a deviation from

¹²⁵ Several lawyers have already committed this cardinal sin. *See, e.g.,* Kim Chandler, *Judge Considers Sanctions Against Attorneys in Prison Case for Using AI in Court Filings*, ASSOCIATED PRESS (May 21, 2025), <https://apnews.com/article/alabama-prisons-ai-8cbaf729dafc2b56bee59545391707c0> [perma.cc/4H5K-6K3F] (discussing a court hearing for a lawyer who used AI to write a brief, that included case citations that do not exist); *Judge Fines Three Lawyers for Citing AI-Generated Fake Cases*, JUSTIA LEGAL NEWS (Feb. 26, 2025), <https://news.justia.com/judge-fines-three-lawyers-for-citing-ai-generated-fake-cases/> [perma.cc/XV4R-64MU] (discussing how three lawyers were sanctioned for citing fake court cases that AI generated in a legal action); Blake Brittain, *Anthropic's Lawyers Take Blame for AI 'Hallucination' in Music Publishers' Lawsuit*, REUTERS, <https://www.reuters.com/legal/legalindustry/anthropics-lawyers-take-blame-ai-hallucination-music-publishers-lawsuit-2025-05-15/> [perma.cc/9XZX-44AL] (May 15, 2025) (discussing an instance of AI creating fake citations that a lawyer then used in a lawsuit).

¹²⁶ *See* Chandler, *supra* note 125; JUSTIA LEGAL NEWS, *supra* note 125; Brittain, *supra* note 125.

¹²⁷ *See, e.g.,* Sarah Schwartz, *AI Gets Math Wrong Sometimes. How Teachers Deal with Its Shortcomings*, EDUC. WK. (Sep. 20, 2024), <https://www.edweek.org/teaching-learning/ai-gets-math-wrong-sometimes-how-teachers-deal-with-its-shortcomings/2024/09> [perma.cc/QS5T-QBXV] (highlighting the frequency in which AI incorrectly answers math problems).

¹²⁸ *See* Kevin Roose, *Can A.I. Be Blamed for a Teen's Suicide?*, N.Y. TIMES, <https://www.nytimes.com/2024/10/23/technology/characterai-lawsuit-teen-suicide.html> [perma.cc/RRG3-N7CD] (Oct. 24, 2024) (detailing a tragic teen suicide that involved a teen's obsession with an AI chatbot that, among other things, appeared to encourage him to commit suicide).

¹²⁹ Some people will rely on the chatbot anyway, of course. But instructions and warnings can serve to make that reliance unreasonable as a matter of law.

¹³⁰ *See* RESTATEMENT (THIRD) OF TORTS: PRODS. LIAB. § 2(b)–(c) (A.L.I. 1998) (providing the standard for design defects and for defects based on inadequate warnings).

¹³¹ *See id.* § 2(a) (providing the standard for a manufacturing defect). A reminder for clarity: we're talking about defectively manufactured AI, and *not* defectively manufactured machines *controlled* by AI. There is no question that if a driverless car's faulty brakes cause an accident, we have a manufacturing defect on our hands. *See id.* But our focus is on whether the car's AI—the brain controlling the car—can be improperly manufactured.

the AI's design, and (2) showing that the deviation caused the harm.¹³² But how? After all, the "design" comprises a "black-box" neural network full of inscrutable numbers.¹³³ Even if we could show that the neural network of the AI in question deviates from the original, how would we prove that those deviations caused the harm? Remember: the parameters in the neural networks are basically inscrutable. How might one attribute meaning to deviations in those parameters?¹³⁴

Relatedly, AI actors are unlikely to provide overt evidence of an alleged manufacturing defect in the way that other kinds of software might. Recall that neural network-based artificial intelligence is qualitatively different from other forms of programming.¹³⁵ Traditional software experiences an error when the software's code contains mistakes, or when the software encounters an event or occurrence that its code did not anticipate.¹³⁶ Errors frequently cause software to crash, unless its code is written to anticipate and handle those errors.¹³⁷ By contrast, an AI's neural network does not break when it encounters an unexpected input.¹³⁸ It may not handle unorthodox or difficult inputs particularly effectively, but it *will* provide an output.¹³⁹ Importantly, it typically won't crash or produce an error message.¹⁴⁰

¹³² See, e.g., *McCorvey v. Baxter Healthcare Corp.*, 298 F.3d 1253, 1257 (11th Cir. 2002) ("Under Florida law, a strict product liability action requires the plaintiff to prove that (1) a product (2) produced by a manufacturer (3) was defective or created an unreasonably dangerous condition (4) that proximately caused (5) injury.").

¹³³ See Hassija et al., *supra* note 108, at 47.

¹³⁴ Even an inductive-reasoning experiment wouldn't work. Suppose you reconstructed a similar scenario and had the AI models play that scenario out to see whether one does worse than the other. In the case of AI, this would prove nothing. Recall that neural network-based AI systems are *probabilistic*. See generally Henderson, *supra* note 43. Even a non-defective model may make a different—even a harm-causing—decision some small percentage of the time. It might be necessary to run the experiment hundreds or thousands of times in order to conclusively identify *any* behavioral differences between the two AI actors.

¹³⁵ See *supra* Part I.A.

¹³⁶ See Amy J. Ko & Brad A. Myers, *A Framework and Methodology for Studying the Causes of Software Errors in Programming Systems*, 16 J. VISUAL LANGUAGES & COMPUTING 41, 43–44 (2005) (discussing types of computer errors, including runtime failures, runtime faults, and software errors).

¹³⁷ JOSEPH C. GIARRATANO, CLIPS USER'S GUIDE: VERSION 6.30, at 109 (2023) (describing the need for preventing crashes due to illegal input values).

¹³⁸ Cameron Buckner & James Garson, *Connectionism*, STAN. ENCYCLOPEDIA OF PHIL., <https://plato.stanford.edu/entries/connectionism/> [perma.cc/FSZ9-4R5F] (Feb. 19, 2015) ("Neural networks exhibit robust flexibility in the face of the challenges posed by the real world. Noisy input or destruction of units causes . . . degradation of function. The net's response is still appropriate, though somewhat less accurate . . . [N]oise and loss of circuitry in classical computers typically result in catastrophic failure."); see also Cesar Torres-Huitzil & Bernard Girau, *Fault and Error Tolerance in Neural Networks: A Review*, 5 IEEE ACCESS 17322, 17326 (2017) ("Graceful degradation is referred to as a low sensitivity to the occurring faults instead of a complete or catastrophic failure.").

¹³⁹ Buckner & Garson, *supra* note 138; see also Jörg Martin & Clemens Elster, *Detecting Unusual Input to Neural Networks*, 51 APPLIED INTEL. 2198, 2198 (2020) ("Evaluating a neural network on an input that differs markedly from the training data might cause erratic and flawed predictions.").

¹⁴⁰ Buckner & Garson, *supra* note 138.

Leaving aside design and manufacturing defects, an AI actor might also be defective due to inadequate instructions or warnings. This is probably the products-liability hook that AI systems are most vulnerable to, for a few reasons. First, the true capabilities of AI systems are both rapidly expanding *and* difficult to understand. Second, current AI systems occasionally get things wrong, and even hallucinate—often in ways that are nonetheless extremely convincing.¹⁴¹ Taken together, this means that for some kinds of AI, there is a very real risk that people might rely on these models when they shouldn't—and the AI product's provided instructions and warnings might not sufficiently canvas the foreseeable harms. An AI might well be found defective for this reason.¹⁴²

On the one hand, sufficient instructions and warnings may not be all that difficult to include. And where an unwarned-about harm occurs despite otherwise robust labeling, the manufacturer may yet avoid liability on grounds the harm was not foreseeable.¹⁴³ But on the other hand, it may be difficult or impossible to provide warnings and instructions sufficient to render an AI non-defective in the case of some AI harms. Chatbots are a good example: how effective is a warning that a given chatbot sometimes hallucinates, if in the course of using it, the chatbot assures you of its own accuracy?

The upshot of this whole discussion is that products liability will indeed be capable of redressing some AI-caused harms. Though it may have some natural resistance to design- and manufacturing-defect claims, it certainly isn't immune to them—and it appears to be unusually vulnerable to inadequate-instructions-and-warnings claims.

This Part has made the case for how negligence and products liability will handle AI harms, and why they will remain effective in the face of this new technology. But others aren't so convinced. Thanks in large part to the fear that negligence and products liability will create an AI-liability gap, many scholars have urged that all AI harms should be subject to strict liability.¹⁴⁴ In their view, if an AI actor causes harm, someone—its manufacturer, seller, or user—should be found to have committed a tort, full-stop.¹⁴⁵ Others have riffed on this idea,¹⁴⁶ or proposed altogether different solutions to the alleged prob-

¹⁴¹ This is especially relevant for large language models designed to output text, images, and audio. *See, e.g.,* Chandler, *supra* note 125 (“[A] firm partner . . . used ChatGPT to research supporting case law but did not verify the information . . . Those citations turned out to be ‘hallucinations’ . . .”); Brittain, *supra* note 125 (“[T]he expert had relied on a legitimate academic journal article, but [the expert] created a citation for it using Anthropic’s chatbot Claude, which made up a fake title and authors . . .”).

¹⁴² *See* RESTATEMENT (THIRD) OF TORTS: PRODS. LIAB. § 2(c) (A.L.I. 1998) (explaining a product can be defective due to insufficient instructions or warnings).

¹⁴³ *See id.* (including foreseeability as a factor in determining whether a product is defective due to inadequate warnings).

¹⁴⁴ *See infra* Part II and II.A for citations to this body of scholarship.

¹⁴⁵ *See infra* Part II and II.A for citations to this body of scholarship.

¹⁴⁶ *See infra* Part II and II.A for citations to this body of scholarship.

lem.¹⁴⁷ I discuss some of these proposals next—and argue that advocating strict liability as a solution to AI harms is doctrinally baseless, and as a normative matter amounts to trying to pass off a battering ram as a skeleton key.

II. PROPOSED SOLUTIONS AND THEIR SHORTCOMINGS

Let's restate the supposed problem: tort's status quo will struggle to redress many harms caused by AI actors. On this and other bases, scholars have suggested a range of innovations. The most vocal cohort proposes strict liability.¹⁴⁸ Others propose variants that contain elements of strict liability—things like presumed negligence with safe harbors,¹⁴⁹ or regulation-derived liability tiers.¹⁵⁰ Another cohort offers a range of more creative solutions not rooted in strict liability,¹⁵¹ with a subset of these focused on particular instances of AI, like driverless vehicles and healthcare-related AI.¹⁵² Notwithstanding my view that the current tort system will work just fine for AI harms, some of these

¹⁴⁷ See *infra* Part II and II.B.

¹⁴⁸ See, e.g., Vladeck, *supra* note 15, at 146 (arguing for strict liability for AI harms); Anat Lior, *AI Strict Liability Vis-à-Vis AI Monopolization*, 22 COLUM. SCI. & TECH. L. REV. 90, 96 (2020) (same); Cristina Carmody Tilley, *Just Strict Liability*, 43 CARDOZO L. REV. 2317, 2377 (2022) (same); Stefan Heiss, *Towards Optimal Liability for Artificial Intelligence: Lessons from the European Union's Proposals of 2020*, 12 HASTINGS SCI. & TECH. L.J. 186, 217–18 (2021) (same, but with damages paid to the state).

¹⁴⁹ See Omri Rachum-Twaig, *Whose Robot Is It Anyway?: Liability for Artificial-Intelligence-Based Robots*, 2020 U. ILL. L. REV. 1141, 1168 (“[E]nhancing an existing liability rule, namely negligence, with supplementary rules that will set a predetermined acceptable level of care applicable to designers and operators of AI-based robots (regardless of whether AI is embedded in the product sold to the consumer or AI capabilities are delivered as a service).”).

¹⁵⁰ See Matthew U. Scherer, *Regulating Artificial Intelligence Systems: Risks, Challenges, Competencies, and Strategies*, 29 HARV. J.L. & TECH. 353, 357 (2016) (proposing different liability standards—strict liability vs. negligence—for certified and non-certified AI systems).

¹⁵¹ See, e.g., Amy L. Stein, *Assuming the Risks of Artificial Intelligence*, 102 B.U. L. REV. 979, 1022–34 (2022) (thoughtfully proposing a reenvisioning of the assumption-of-risk doctrine in light of the benefits and risks of AI); Jin Yoshikawa, *Sharing the Costs of Artificial Intelligence: Universal No-Fault Social Insurance for Personal Injuries*, 21 VAND. J. ENT. & TECH. L. 1155, 1178 (2019) (proposing chucking out most of tort law in favor of a no-fault liability insurance system modeled after New Zealand's social-insurance scheme).

¹⁵² See, e.g., A. Michael Froomkin, Ian Kerr & Joelle Pineau, *When AIs Outperform Doctors: Confronting the Challenges of a Tort-Induced Over-Reliance on Machine Learning*, 61 ARIZ. L. REV. 33, 81–83 (2019) (addressing the interesting challenge of machine-learning tools surpassing physicians in diagnostics, the attendant downstream unintended consequences, and ways of adjusting the duty of care to continue to ensure a human in the loop); Kevin M.K. Fodouop, Note, *The Road to Optimal Safety: Crash-Adaptive Regulation of Autonomous Vehicles at the National Highway Traffic Safety Administration*, 98 N.Y.U. L. REV. 1358, 1387–88 (2023) (proposing that regulators “follow the [autonomous vehicle] industry’s lead in using data tracking and simulation as post-crash information-forcing tools” to develop “crash-adaptive optimal regulation”); Jeffrey K. Gurney, *Sue My Car Not Me: Products Liability and Accidents Involving Autonomous Vehicles*, 2013 U. ILL. J.L. TECH. & POL’Y 247, 274–77 (proposing shifting liability depending on what kind of driver is deploying the autonomous vehicle).

proposals—especially those promoting private insurance mechanisms—contain meritorious elements.¹⁵³

But others do not—categorical strict liability least of all. Section A of this part discusses arguments for strict liability’s applicability to AI harms as well as its potential shortcomings.¹⁵⁴ Section B then discusses other proposed frameworks for addressing AI harms.¹⁵⁵

A. Strict Liability

Strict liability assigns liability irrespective of care, precaution, or blameworthiness.¹⁵⁶ The most careful lion-tamer is just as liable as the least careful one¹⁵⁷; the scrupulous dynamite engineer is no less on the hook than a blundering backyard blaster.¹⁵⁸ These examples hint at strict liability’s most common applications: assigning tort liability for harms caused by abnormally dangerous activities,¹⁵⁹ or through the handling of certain kinds of animals.¹⁶⁰ There are a few other common-law applications for strict liability, too—but only a few. One is for manufacturing defects in the products liability context.¹⁶¹ Another is certain instances of private and public nuisance.¹⁶² Apart from these, lawmak-

¹⁵³ For example, Amy Stein’s suggestion that assumption of risk will or ought to play an important role in AI harms is spot on; and though Jin Yoshikawa’s proposal for trading our tort regime for social insurance goes too far, I share his view that insurance is a useful mechanism for redressing many AI harms. See generally Stein, *supra* note 151; Yoshikawa, *supra* note 151.

¹⁵⁴ See *infra* Part II.A.

¹⁵⁵ See *infra* Part II.B.

¹⁵⁶ Richard A. Epstein, *A Theory of Strict Liability*, 2 J. LEGAL STUD. 151, 152 (1973).

¹⁵⁷ See *Union Pac. R.R. Co. v. Nami*, 498 S.W.3d 890, 896–97 (Tex. 2016) (“[A] person who owns, possesses, or harbors a wild animal is strictly liable for its actions.”).

¹⁵⁸ See *Klein v. Pyrodyne Corp.*, 810 P.2d 917, 922 (Wash. 1991) (noting dynamite is subject to strict liability, and holding that in Washington, so are fireworks).

¹⁵⁹ See RESTATEMENT (SECOND) OF TORTS § 520 (A.L.I. 1977); see also *Ind. Harbor Belt R.R. Co. v. Am. Cyanamid Co.*, 916 F.2d 1174, 1176 (7th Cir. 1990) (“The key provision is section 520, which sets forth six factors to be considered in deciding whether an activity is abnormally dangerous and the actor therefore strictly liable.”).

¹⁶⁰ See, e.g., *Cook v. Whitsell-Sherman*, 796 N.E.2d 271, 276 (Ind. 2003) (“[P]ossession of a wild animal is, like blasting, an unreasonably dangerous activity subjecting the actor to strict liability.”); *Fletcher v. Richardson*, 603 S.W.2d 734, 735 (Tenn. 1980) (stating an owner who knowingly keeps a vicious animal has “absolute liability” for harms it causes).

¹⁶¹ RESTATEMENT (THIRD) OF TORTS: PRODS. LIAB. § 2(a) (A.L.I. 1998).

¹⁶² See *Cox v. City of Dallas*, 256 F.3d 281, 290 (5th Cir. 2001) (“The rules of strict liability, i.e., liability imposed without regard to the defendant’s negligence or intent to harm, are frequently applied to abnormally dangerous activities, see RESTATEMENT (SECOND) OF TORTS § 519 (1977), although they are imposed in other nuisance situations as well.” (footnote omitted)).

ers have occasionally imposed strict liability by statute¹⁶³—but rarely, and often only to codify the common-law rules.¹⁶⁴

In short, strict liability is and has been for nearly two centuries the exception, not the rule.¹⁶⁵ But despite this, many scholars argue today that strict liability ought to govern virtually all harms caused by AI actors.¹⁶⁶ This proposal is stunning not only for its breadth, but also because it is unmoored from any of the prevailing theories of tort liability. Let's start with a look at the proposals themselves.

One early and recurring argument for strict liability is that products liability might fail to reach an AI harm that we would have wanted it to redress.¹⁶⁷ David Vladeck made this argument earlier than most: in 2014 he proposed strict liability for harms caused by AI actors.¹⁶⁸ Vladeck presciently noted that, at least in the context of driverless vehicles, strict liability couldn't be based on the traditional "ultra-hazardous" rationale, because driverless vehicles were likely to be comparatively safe technologies.¹⁶⁹ Accordingly, he called for a "true strict liability regime," based on four rationales: the importance of "redress for persons injured through no fault of their own"; the societal value of burden-sharing; transaction-cost minimization; and innovation promotion on account of liability predictability.¹⁷⁰

Others writing more recently have reached the same shore charting different routes. For example, Anat Lior has proposed that we analogize AI actors to AI agent-servants, and meld respondeat superior with strict liability to hold a principal strictly liable.¹⁷¹ Lior locates additional support for strict liability in George Fletcher's "nonreciprocal harms" framework.¹⁷² According to Lior, AI

¹⁶³ See, e.g., Comprehensive Environmental Response, Compensation, and Liability Act, 42 U.S.C. § 9607(a); *United States v. Monsanto Co.*, 858 F.2d 160, 167 (4th Cir. 1988) (interpreting 42 U.S.C. § 9607(a) stating "we agree with the overwhelming body of precedent that has interpreted section 107(a) as establishing a strict liability scheme").

¹⁶⁴ For example, many states have dog-bite statutes that impose strict liability. See, e.g., CAL. CIV. CODE § 3342 (West 2025). But some laws imposing strict liability do go beyond mere codification of the common law. For example, the Comprehensive Environmental Response, Compensation, and Liability Act "is a strict liability statute; its language sets forth categories of responsible persons who can be held liable for cleanup costs even where they have not acted intentionally or negligently in creating the harm." Lynda J. Oswald, *Strict Liability of Individuals Under CERCLA: A Normative Analysis*, 20 B.C. ENV'T AFFS. L. REV. 579, 583 (1993) (footnote omitted).

¹⁶⁵ See WILLIAM M. LANDES & RICHARD A. POSNER, *THE ECONOMIC STRUCTURE OF TORT LAW* 3 (1987) ("By the middle of the nineteenth century it was settled in both England and America that the normal standard of liability in accident cases was negligence; only exceptionally was it strict liability.").

¹⁶⁶ See *infra* notes 167–179 and accompanying text.

¹⁶⁷ Vladeck, *supra* note 15, at 146.

¹⁶⁸ *Id.*

¹⁶⁹ *Id.*

¹⁷⁰ *Id.* at 146–47.

¹⁷¹ Lior, *supra* note 61, at 1084 ("A strict liability regime aligned with the *respondeat superior* doctrine will enable this to happen.").

¹⁷² Lior, *supra* note 148, at 95–96.

actors can inflict greater and more frequent harm on their environments than can be inflicted on them—an imbalance that, based on her read of Fletcher, justifies strict liability.¹⁷³

Cristina Carmody Tilley suggests that Fletcher’s reciprocity theory is too vague to support strict liability on moral grounds,¹⁷⁴ but she comes armed with her own justice-oriented arguments. Tilley stresses the “relational limits” of AI,¹⁷⁵ building on one of her earlier arguments that strict liability ought to be “theorized as an unapologetically just sector of tort.”¹⁷⁶ According to Tilley, “[c]omplete delegations to AI produce relationships that are care-resistant, trilateral, extraordinarily dangerous, and properly subject to a strict liability rule.”¹⁷⁷

Finally, though focused on mechanized and routinized processes and not AI specifically, Lee Fennell makes perhaps the most persuasive argument for strict liability. Building on earlier work by Robert Cooter and Ariel Porat,¹⁷⁸ Fennell proposes that machine strict liability is the best way to cut through the following problematic distortion caused by the disparate treatment of machine vs. human harms: negligence tends to judge machines in the aggregate, but humans in the moment.¹⁷⁹ I respond to Fennell more fully in the next subpart.

Still others riff on strict liability in various ways and to varying degrees. Stefan Heiss favors strict liability, but with payments made to the state.¹⁸⁰ Matthew Scherer proposes a regulatory regime in which AI systems that undergo a certification process would enjoy a negligence standard, whereas uncertified AI systems—the unwashed masses—would face strict liability.¹⁸¹ And Omri Rachum-Twaig proposes a model of presumed negligence,¹⁸² which reads, to me, like a diet version of strict liability.

In my view, none of these proposals are persuasive; categorical strict liability for AI harms is neither normatively warranted, nor doctrinally grounded. In the next few subparts, I explain why.¹⁸³

¹⁷³ *Id.* at 97.

¹⁷⁴ Tilley, *supra* note 148, at 2332.

¹⁷⁵ *Id.* at 2377.

¹⁷⁶ *Id.*

¹⁷⁷ *Id.* (emphasis omitted).

¹⁷⁸ See generally Robert Cooter & Ariel Porat, *Lapses of Attention in Medical Malpractice and Road Accidents*, 15 THEORETICAL INQUIRIES L. 329 (2014) (offering potential solutions to inefficiencies in legal standards for lapses in precaution).

¹⁷⁹ Lee Anne Fennell, *Accidents and Aggregates*, 59 WM. & MARY L. REV. 2371, 2384–87 (2018).

¹⁸⁰ Heiss, *supra* note 148, at 186, 217–18.

¹⁸¹ Scherer, *supra* note 150, at 357.

¹⁸² Rachum-Twaig, *supra* note 149, at 1167–69.

¹⁸³ See *infra* Part II.A.1–Part II.A.2.

1. Normative Shortcomings of Strict Liability for AI Harms

To begin with, strict liability isn't merely an incentive to act differently; it is an incentive not to act at all.¹⁸⁴ Use dynamite to demolish that downtown building if you must—but *only* if you must. You love and want to keep your dog despite its known tendency to bite—so keep it muzzled, or else. Your company can't be competitive without outsourcing manufacturing to unregulatable, hard-to-sue overseas factories—so if they screw something up and hurt someone, that's on you.¹⁸⁵ Strict liability is a heavy-handed disincentive that jacks up the cost of dangerous or deeply undesirable activities, sometimes to the point of driving those activities to extinction.¹⁸⁶

Plainly stated, strict liability simply doesn't fit the bill. Despite considerable handwringing, there is no evidence that AI activity, *as a category*, is abnormally dangerous, or socially undesirable. For starters, talking about "AI activity" is almost as unwieldy as talking about "human activity"; as a category, it is impossibly broad. Today AI helps us discover new drugs,¹⁸⁷ drives us to work,¹⁸⁸ and makes a perfect cappuccino.¹⁸⁹ The expanding list of its capabilities is already nearly endless—and most of what AI actors can do wouldn't qualify as abnormally dangerous if a human were doing it.¹⁹⁰ What's more, all early evidence suggests that for activities that pose normal dangers—things

¹⁸⁴ *Ind. Harbor Belt R.R. Co. v. Am. Cyanamid Co.*, 916 F.2d 1174, 1177 (7th Cir. 1990) ("By making the actor strictly liable . . . we give him an incentive, missing in a negligence regime, to experiment with methods of preventing accidents that involve not greater exertions of care, assumed to be futile, but instead . . . changing, or reducing (perhaps to the vanishing point) the activity giving rise to the accident.").

¹⁸⁵ See Keith N. Hylton, *The Law and Economics of Products Liability*, 88 NOTRE DAME L. REV. 2457, 2460 (2013) ("[S]trict liability enables the market to distinguish (and ultimately usher out of circulation) products generated from low-quality manufacturing processes.").

¹⁸⁶ *Ind. Harbor Belt R.R. Co.*, 916 F.2d at 1177; see also LANDES & POSNER, *supra* note 165, at 69 ("Strict liability gives the potential injurer an incentive to reduce his activity when such an adjustment is less costly than his taking care (or letting the accident occur and paying the victim's damages), because he bears the costs of an accident that could be avoided by such an adjustment.").

¹⁸⁷ Yuxi Wang, Zelin Hu, Junbiao Chang & Bin Yu, *Thinking on the Use of Artificial Intelligence in Drug Discovery*, 68 J. MED. CHEMISTRY 4996, 4996 (2025), <https://doi.org/10.1021/acs.jmedchem.5c00373>.

¹⁸⁸ See AUTONOMOUS VEHICLE INDUS. ASS'N, *supra* note 82.

¹⁸⁹ See Lauren Martinez, *This Highly Capable AI-powered Robot Barista Is Set to Debut at a San Jose Café*, ABC7 NEWS (Apr. 17, 2025), <https://abc7news.com/post/ai-powered-robot-barista-richtech-robotics-is-set-debut-ncm-caf-san-jose/16188984/> [perma.cc/UU64-R2SK].

¹⁹⁰ No doubt AI will eventually be deployed to assist with abnormally dangerous activities as well—things like handling explosives. And why shouldn't it? Such activities are dangerous for the explosive handler, too—and human errors can have devastating consequences. AI will undoubtedly make many abnormally dangerous activities safer, both by reducing the likelihood of error, and by removing humans from the zone of danger. There is no reason why strict liability shouldn't continue to apply in these contexts.

like driving—AI actors can perform them more safely than we can.¹⁹¹ This is a good reason to *incentivize* AI development and deployment. But across-the-board strict liability would do the opposite.¹⁹²

On this point, some scholars have suggested that strict liability would actually spur innovation, on predictability grounds.¹⁹³ The argument goes something like this: strict liability creates certainty, and certainty is good for innovation; therefore, strict liability is good for innovation.¹⁹⁴ This is a common form of flawed reasoning called the fallacy of the undistributed middle—basically, a non sequitur.¹⁹⁵ Consider the following logically identical—and obviously flawed—syllogism. Getting fired gives people more free time, and free time is good for mental health; therefore, getting fired is good for mental health.

It's true that, *in a vacuum*, predictability supports innovation. But predictability is innovation-promoting because it helps illuminate the profitable path forward. If the road revealed turns out to be full of unavoidable landmines, as would likely be the case with strict liability, then predictability incentivizes nothing—except, perhaps, avoiding the road entirely.

Furthermore, these and other scholars fail to address (adequately or at all) the liability asymmetry that strict liability for AI harms would engender.¹⁹⁶ Suppose an AI actor is strictly liable, whereas a comparable human actor is liable only for negligence. If we assume similar rates and magnitudes of harm, liability for the AI actor's actions will be much more expensive than the human's liability, because claims will succeed more often under strict liability than negligence.¹⁹⁷ The human's actions will achieve liability-cost parity with the AI actor only if he is considerably more harm-causing—in terms of frequency, magnitude, or both—than his AI counterpart.

¹⁹¹ See AUTONOMOUS VEHICLE INDUS. ASS'N, *supra* note 82 (highlighting that autonomous vehicles are involved in less crashes than human drivers).

¹⁹² See, e.g., Hylton, *supra* note 185, at 2479 (“[S]trict liability internalizes the injury risks introduced by the new product. The end result is that some of the benefits of innovation are externalized to consumers . . . while the new risk costs are fully internalized to the producer by liability This implies that innovation incentives are not optimal under strict liability.”).

¹⁹³ See Vladeck, *supra* note 15, at 147 (explaining strict liability “may better spur innovation than a less predictable system”); see also Lior, *supra* note 148, at 94 (“A strict liability regime may even encourage innovation, if it provides manufacturers with a certain, predictable legal framework to operate in.”).

¹⁹⁴ See Vladeck, *supra* note 15, at 147 (highlighting that a stable system more sufficiently encourages responsibly innovation).

¹⁹⁵ See Richard Nordquist, *Undistributed Middle (Fallacy)*, THOUGHTCO., <https://www.thoughtco.com/undistributed-middle-fallacy-1692453> [perma.cc/8D83-TGDY] (May 11, 2025) (“The undistributed middle is a logical fallacy of deduction in which the middle term of a syllogism is not distributed in at least one of the premises.” (emphasis omitted)).

¹⁹⁶ See Vladeck, *supra* note 15, at 147; Lior, *supra* note 148, at 94.

¹⁹⁷ See LANDES & POSNER, *supra* note 165, at 65 (“Under strict liability a claim arises every time there is an accident caused by someone worth suing; under negligence there is a claim only if the victim thinks that he can show that the defendant failed to use due care.”).

Simple microeconomics describes strict liability's disincentive effects on innovation just as readily. Vladeck lauds the cost-sharing effects of strict liability,¹⁹⁸ but as a practical matter, "cost-sharing" is just a euphemism for "costs more." AI firms will pass the cost of strict liability on to their customers through price increases. When costs go up, demand goes down—and demand is the whole innovation ballgame.¹⁹⁹

There are a few reasonable counters to these economic arguments against strict liability. One argument is that the incentives to develop AI might still outweigh the disincentives. A fair point—but true of any number of things we might apply strict liability to, but don't. For example, if we held car manufacturers strictly liable not only for manufacturing defects but for design defects as well, they would still make cars, and we would still buy them. But they would be incentivized to innovate less, and to focus more on fine-tuning the safety of tried-and-true models. As for us consumers, we would end up paying more for fewer options. Likewise for AI. Under strict liability, we would no doubt continue to see innovation—but at a slower pace, and of a different kind, at a higher cost.

Another more persuasive counter—this one rejoining the "liability asymmetry" argument above—comes from Lee Fennell's work on accidents, aggregates, and the insight that "[t]ort law deals in lumps."²⁰⁰ Fennell homes in on the fact that human-caused harms are evaluated on a per-event basis, whereas machine-caused harms are often evaluated in the aggregate, with a focus on error rate.²⁰¹ What's more, some human-caused harms are the result of unavoidable human lapses—and so negligence occasionally "hold[s] people responsible for shortcomings they could not prevent"²⁰² Not so with machines. So long as there is no cost-effective means of improving a machine's error rate, its deployment might well be deemed non-negligent.²⁰³ Thus, for Fennell, comparing humans to machines is not "apples to apples"—it's "apples

¹⁹⁸ See Vladeck, *supra* note 15, at 146 ("[A] strict liability regime is warranted because, in contrast to the injured party, the vehicle's creators are in a position to either absorb the costs, or through pricing decisions, to spread the burden of loss widely.").

¹⁹⁹ Fulya Taşel, Ebru Bayarçelik & Sinan Apak, *Innovation Decision of Export Companies and Effect on Firms Financial Performance*, 3 FIN. & CREDIT ACTIVITY PROBS. THEORY & PRAC. 176, 182 (2019) ("[R]esults show that . . . market demand is the most important factor influencing the innovation decision."); Jacob Schmookler, *Economic Sources of Inventive Activity*, 22 J. ECON. HIST. 1, 1 (1962) ("The fundamental conclusion of this paper is that technological progress is intimately dependent on economic phenomena New goods and new techniques are unlikely to appear, and to enter the life of society without a pre-existing,—albeit possibly only latent—demand.").

²⁰⁰ Fennell, *supra* note 179, at 2373.

²⁰¹ *Id.* at 2394–95.

²⁰² *Id.* at 2391–92.

²⁰³ *Id.* at 2395.

to orchards.”²⁰⁴ On these grounds, she proposes that strict liability for machine-caused harms would “merely level[] the playing field.”²⁰⁵

Still, this position does not solve the liability-cost-parity problem. True, it addresses a notable incongruity that risks unfairness, especially in the absence of effective insurance: negligence treats harms caused by unavoidable human lapses differently than harms caused by unavoidable machine errors. But the liability-cost burdens that strict liability imposes on innovation remain fully intact.

A notable feature of Fennell’s argument is that it indirectly illuminates a potential employment concern.²⁰⁶ All else being equal, lapse-prone humans carry greater liability costs—and may therefore be less employable—than their error-prone-robot counterparts. Accordingly, strict liability for AI harms would shift the balance in the human employee’s favor. But even stipulating this as undesirable, I see no principled basis for using tort law as a vehicle for implementing policy targeted at employment protection. Insofar as our lumpy tort system works an unfairness on alleged human tortfeasors, it would be better, in my view, to deal with that issue head on. For example, by implementing some version of Cooter and Porat’s proposed lapse defense.²⁰⁷ Furthermore it would be difficult to anchor an employment-protection justification for strict liability in our prevailing theories of tort law.²⁰⁸

There *is* a version of strict liability that might spur innovation: strict liability that compensates only for medical expenses and lost wages (or some other limitation that significantly caps recovery). As it turns out, this roughly describes the bargain we struck decades ago when we passed the first workers’ compensation statutes.²⁰⁹ But applying such a system to so broad a category of harms would be an even greater detour from American tort law than imposing categorical simple strict liability. (Workers’ compensation is far from uncontroversial,²¹⁰ and its exclusive-remedy requirement was and is a big deal—so

²⁰⁴ *Id.*

²⁰⁵ *Id.* at 2395–96 (“Given the tendency toward evaluative disaggregation in the human case and the tendency toward evaluative aggregation in the mechanized case, a (nominal) negligence standard in the former paired with strict liability in the latter merely levels the playing field. And it does so in a way that conserves on information costs because it eliminates the need to investigate the human’s overall ‘error rate.’”).

²⁰⁶ See *id.* (providing Fennell’s argument that machines be governed by strict liability and humans by negligence).

²⁰⁷ See Cooter & Porat, *supra* note 178, at 339–58. Cooter and Porat focus “on the question of allowing or disallowing a second-order reasonableness defense for lapses.” *Id.* at 331. They propose that “[a]llowing such a defense would eliminate inefficiencies which result from liability for unreasonable first-order but reasonable second-order behavior . . .” *Id.*

²⁰⁸ See *infra* Part III.

²⁰⁹ MARGARET C. JASPER, WORKERS’ COMPENSATION LAW 11 (2d ed., 2008) (noting that typically, “[w]orkers’ compensation is the exclusive remedy for an employee to obtain wage replacement for work-related injuries or illnesses, and your employer is generally immune from liability”).

²¹⁰ See generally Richard A. Epstein, *The Historical Origins and Economic Structure of Workers’ Compensation Law*, 16 GA. L. REV. 775 (1982) (discussing historical debates surrounding employer

much so that we still refer to it as the “great trade-off.”)²¹¹ One version of this might be to limit compensation only to reimbursement of medical expenses directly related to the AI-caused harm. This approach would ensure that the most worrying class of AI harms—physical injury—is redressable, without opening AI companies up to the full weight of litigation that the plaintiff’s bar could otherwise impose on them.

Of course, this approach would incentivize innovation only if it were exclusive, rather than additive of other forms of liability. In other words, it would need to be its own “Great Trade-Off”: increased frequency of liability for AI harms²¹² in exchange for significantly decreased magnitude of liability.

This approach is beyond improbable—not to mention a bad idea. The sheer magnitude of harms it would cover is immense. (In the next subpart I provide some examples that show just *how* broadly any categorical rule would sweep.²¹³) And state legislatures the country over would be unlikely to create what would amount to parallel tort regimes just for AI harms. What’s more, it would create the incentive to stuff AI into virtually everything. After all, under this rule, if a defendant can blame AI, then they can cap the damages. This workers’ compensation-like approach *might* work in, for example, the U.K., where tort law already is less punitive, and where theories of compensation are more narrowly defined.²¹⁴ But even that would be a stretch.

This brings us to the second half of this discussion: foisting strict liability categorically onto AI harms is doctrinally unsound as well.

2. Doctrinal Shortcomings of Strict Liability for AI Harms

Leaving aside the many normative shortcomings of using strict liability to redress AI harms, doing so also has little support or justification in pre-existing doctrine.

The cornerstone of tort law in the United States is fault, which we determine principally through negligence.²¹⁵ Where a product is involved, we often

liability); Martha T. McCluskey, *The Illusion of Efficiency in Workers Compensation Reform*, 50 RUTGERS L. REV. 657 (1998) (challenging recent reforms in workers’ compensation).

²¹¹ Lloyd Harger, *Workers’ Compensation, A Brief History*, FLA. DEP’T OF FIN. SERV., <https://my.floridacfo.com/division/wc/infofaqs/history> [perma.cc/A4SN-PZK8].

²¹² See LANDES & POSNER, *supra* note 165, at 65 (“Under strict liability a claim arises every time there is an accident caused by someone worth suing; under negligence there is a claim only if the victim thinks that he can show that the defendant failed to use due care.”).

²¹³ See *infra* Part II.A.2.

²¹⁴ Marco Cappelletti, *Comparative Reflections on Punishment in Tort Law* (stating that the primary aim of tort law in the U.K. is to compensate claimants, rather than punish the defendant, with damages only rewarded in certain cases), in *FRENCH CIVIL LIABILITY IN COMPARATIVE PERSPECTIVE* 329, 337 (Jean-Sébastien Borghetti & Simon Whittaker eds., 2019).

²¹⁵ *Ind. Harbor Belt R.R. Co. v. Am. Cyanamid Co.*, 916 F.2d 1174, 1177 (7th Cir. 1990) (“The baseline common law regime of tort liability is negligence.”).

turn to products liability, which also stresses reasonableness, except in the narrow case of manufacturing defects.²¹⁶ For the foreseeable future, most AI actors will probably (but far from certainly) be found to fit the definition of a product,²¹⁷ so all these strict-liability proposals amount to one of two things. Either add a new category of products liability for AI harms, or alternatively, pull AI harms out of the products-liability framework, and force strict liability some other way.

What would a principled augmentation of products liability look like? Most arguments favoring strict liability point to how products liability will struggle to reach AI harms²¹⁸—a suggestion this article analyzed in detail in Part I and mostly rejected.²¹⁹ But products liability struggles to reach plenty of non-AI harms, too. If I fall off a ladder or cut myself with my chef's knife, products liability won't help me; the risks were open and obvious (ladders are tall; knives are sharp). I might have a heart attack while running on my treadmill, but no one is going to wake up NordicTrack's general counsel in the middle of the night if I do (NordicTrack warns me that I ought to talk with my doctor before exercising strenuously²²⁰). Injuries like these happen a hundred thousand times a day—yet we don't see a need to make ladders, knives, or treadmills subject to strict liability.

But suppose we decide AI is fundamentally different, and we augment products liability, such that harms caused by AI products *are* subject to blanket strict liability. Further suppose that my ladder is equipped with AI that beeps if it detects that it has become unstable. The beep startles me, and I fall off, breaking my leg. Strict liability. Or perhaps the AI in my chef's knife detects that it's struggling to cut through carrots. The AI knife auto-sharpens itself in the holding block next time I put it away. Because it's now razor-sharp, I cut the tip of my finger off, though a duller knife would have stopped at the fingernail. Strict liability. Or my AI treadmill, detecting that I'm not sweating despite having selected a strenuous workout, kicks the speed up a few notches. As a result of my increased heart rate, I have a heart attack. Strict liability.

²¹⁶ See RESTATEMENT (THIRD) OF TORTS: PRODS. LIAB. § 2(a) (A.L.I. 1998).

²¹⁷ See *id.* § 19(a) (“A product is tangible personal property distributed commercially for use or consumption. Other items, such as real property and electricity, are products when the context of their distribution and use is sufficiently analogous to the distribution and use of tangible personal property that it is appropriate to apply the rules stated in this Restatement.”); *c.f.* James M. Beck, *Guest Post—Is Artificial Intelligence a “Product”? The Third Circuit Says, “No.”*, LEXOLOGY (Mar. 27, 2020), <https://www.lexology.com/library/detail.aspx?g=ab772071-757a-456a-a027-8ac4428c78b3> [perma.cc/Q5MU-WUZ5] (discussing a court case that held an AI algorithm does not qualify as a “product” under the Restatement).

²¹⁸ Lior, *supra* note 171, at 1057–58.

²¹⁹ See *supra* Part I.

²²⁰ See *Frequently Asked Questions: Treadmills*, NORDICTRACK BLOG, <https://www.nordictrack.com/learn/faq-treadmills> [perma.cc/PT5D-LAHY] (Aug. 29, 2025) (providing a disclaimer that users consult their doctors before altering their fitness routines).

None of this would serve the purpose of products liability; none of the above products are defective.²²¹ Yet this is what subjecting AI products to strict liability could result in—AI products that soon will be utterly ubiquitous.

Leaving products liability aside, these examples also help explain why it would not make sense to shoehorn AI actors into the “abnormally dangerous” framework alongside dynamite or pet wolves.²²² No one would suggest the AI actors built into the ladder, knife, or treadmill engage in abnormally dangerous activities. The restatements would agree. Of the second restatement’s six factors for “determining whether an activity is abnormally dangerous,” these three examples would satisfy no more than two, and as few as none.²²³ Nor would these activities meet the third restatement’s simplified definition of abnormally dangerous activity, which requires that the activities not only pose a “highly significant risk of physical harm,” but that “the activity is not one of common usage.”²²⁴

It’s worth noting that the “abnormally dangerous” bucket represents the only real possibility of foisting strict liability on AI actors through the common law. All other approaches would likely require legislation. And if any state legislature actually passed such a statute, it would immediately create one of the most polarized legal patchworks of the modern era.²²⁵ What’s more, given the growing ubiquity and necessity of AI in all economic sectors, it would not surprise me if such a statute were challenged successfully on Dormant Commerce Clause grounds.²²⁶ Of course, the federal legislature has the power to make strict liability for AI harms the law of the land. But as a practical matter, preemptive federal AI tort legislation seems even less likely than new consensus among the several states.

So categorical strict liability is a poor fit, both in practical terms and as a matter of doctrine. What else is on the table?

²²¹ See RESTATEMENT (THIRD) OF TORTS: PRODS. LIAB. § 1 (A.L.I. 1998).

²²² See RESTATEMENT (SECOND) OF TORTS § 520 (A.L.I. 1977) (providing the factors to consider in deciding whether an activity is “abnormally dangerous”).

²²³ See *id.* Only (a) and (b)—“existence of a high degree of risk” and “likelihood that the harm that results from it will be great”—might arguably apply. See *id.* (a), (b).

²²⁴ RESTATEMENT (THIRD) OF TORTS: LIAB. FOR PHYSICAL AND EMOTIONAL HARM § 20 (A.L.I. 2010). Notably, this latter requirement is not without its critics. See, e.g., Steven Shavell, *The Mistaken Restriction of Strict Liability to Uncommon Activities*, 10 J. LEGAL ANALYSIS 1, 1 (2018) (arguing for the elimination of the uncommon activity requirement).

²²⁵ This is a normative consideration, but I place it here because it flows directly from the point that “abnormally dangerous” is the only plausible common-law hook for a strict-liability-for-AI rule. See RESTATEMENT (SECOND) OF TORTS § 520.

²²⁶ Note, *The Dormant Commerce Clause and Moral Complicity in a National Marketplace*, 137 HARV. L. REV. 980, 980 (2024) (“Recent litigation and scholarship have invoked the dormant commerce clause as a limit on state ethical innovation, arguing that state regulations that impose unique compliance costs may need to yield to a uniform national market . . .”).

B. Other Proposed Solutions

Besides strict liability, scholars have put forward a range of alternative liability schemes for handling AI harms. They run the gamut—from AI personhood to publicly funded social insurance à la New Zealand, and everything in-between.²²⁷ All share a common aim: coping with the supposed problem of AI harms that current tort law may have difficulty redressing.

AI personhood in some form has a surprising number of adherents. Iria Giuffrida asks why not “treat AI systems as we do corporations.”²²⁸ Brian Gannon, Scott Schweikart, and Benedict See either propose or evaluate versions of AI personhood as well.²²⁹ One advantage of doing so, of course, would ostensibly be to make AI actors vulnerable to negligence. But for the reasons we discussed in Part I.A, this probably isn’t sensible until and unless AI actors warrant personhood on sentience grounds.²³⁰ Nor does corporate personhood offer a good analogy. For one thing, the generally accepted bases for corporate personhood don’t apply to AI.²³¹ For another, a corporation always includes human beings that are at least theoretically reachable through the criminal law.²³² It’s hard to imagine a theory of AI personhood that would guarantee a similar level of buck-stops-here accountability.

Another substantial cohort proposes managing AI harms through regulation. For instance, Matthew Scherer proposes the creation of a new regulatory body to oversee the certification of AI systems—with certified and uncertified AI being subject to different liability standards.²³³ A second group focuses on

²²⁷ See *infra* notes 228–246 and accompanying text.

²²⁸ Iria Giuffrida, *Liability for AI Decision-Making: Some Legal and Ethical Considerations*, 88 FORDHAM L. REV. 439, 444 (2019) (“The idea of granting AI systems personhood is not science fiction but rather a legal fiction.”).

²²⁹ See Gannon, *supra* note 70, at 169 (“A way around this complication may be to assign legal personality—and thus liability—directly to the agent”); Scott J. Schweikart, *Who Will Be Liable for Medical Malpractice in the Future? How the Use of Artificial Intelligence in Medicine Will Shape Medical Tort Law*, 22 MINN. J.L. SCI. & TECH. 1, 22 (2021) (“[S]ome mix of the three possibilities outlined above—AI personhood, common enterprise liability, and a new standard of care—is likely to be implemented.”); Benedict See, *Paging Doctor Robot: Medical Artificial Intelligence, Tort Liability, and Why Personhood May Be the Answer*, 87 BROOK. L. REV. 417, 437 (2021) (“[T]he best course of action is adopting personhood for AI coupled with mandatory medical malpractice insurance.”).

²³⁰ See *supra* Part I.A.

²³¹ See, e.g., Elizabeth Pollman, *Corporate Personhood and Limited Sovereignty*, 74 VAND. L. REV. 1727, 1749 (2021) (“[T]he Court has long accorded rights to corporations based on the rationale that corporations represent associations of people from whom such rights are derived or on an instrumental basis to protect the rights of parties outside the corporation.”).

²³² See Robert B. Thompson, *Piercing the Corporate Veil: An Empirical Study*, 76 CORN. L. REV. 1036, 1041 (1991) (describing instances that justify piercing the corporate veil and holding its human shareholders and decision-makers liable).

²³³ See Scherer, *supra* note 150, at 357.

regulating AI in particular industries, such as self-driving cars²³⁴ and health-care.²³⁵ Others train their attention on handling AI harms in the European Union, urging regulatory harmonization across the E.U.²³⁶ In my view regulation for AI harms is inevitable—and we ought to handle it in the same way we regulate human harms: by industry, or by narrower categories of harm. And on this view, our current regulatory regimes likely already contain all the flexibility they need to account for harms that are unique to artificial intelligence.

A few scholars propose augmentations to negligence, or to the interpretive standards employed in negligence suits. Andrew Selbst proposes some of these, asserting that “negligence will need a set of outside interventions to have a real chance at providing redress for harms that result from the use of AI.”²³⁷ Jane Bambauer offers the possibility of a “reasonable AI developer” standard “where errors should be expected to be less frequent than, though qualitatively different from . . . the errors made by a reasonable person.”²³⁸ And Ryan Abbott suggests that the reasonable person standard might organically turn into a kind of “reasonable robot” standard.²³⁹ (Notably this interesting proposal is not aimed at holding AI responsible.) These proposals are interesting and may well be experimented with by courts looking to apply the law of negligence to AI harms. But raising them and seeing what the courts do with them—if anything—is enough. I don’t read any of these proposals as calls for new statutes or regulations.

Among the most interesting proposals is Amy Stein’s suggestion to breathe new life into the assumption-of-risk defense.²⁴⁰ She suggests that informed-consent frameworks might reinvigorate assumption-of-risk doctrine in a way that strikes “a better balance between two of tort law’s competing goals: facilitating innovation and compensating the injured.”²⁴¹ She believes the creation and deployment of an objective standard would enable the assumption-of-risk doctrine to provide good incentives all around—for manufacturer innova-

²³⁴ See Fodouop, *supra* note 152, at 1387–88 (pushing for a regulatory framework surrounding self-driving cars).

²³⁵ See Charlotte A. Tschider, *Medical Device Artificial Intelligence: The New Tort Frontier*, 46 *BYU L. REV.* 1551, 1552 (2021) (“This Author then recommends an alternative conception of the regulatory-tort allocation for AI machines that will create a more comprehensive and complementary safety and compensatory model.”).

²³⁶ See Béatrice Schütte & Lotta Majewski, *Private Liability for AI-Related Harm: Towards More Predictable Rules for the Single Market*, 6 *MKT. & COMPETITION L. REV.* 123, 123 (2022) (“The envisioned level playing field throughout the Single Market justifies harmonisation of many aspects of damages liability for AI-related harm.”).

²³⁷ Andrew D. Selbst, *Negligence and AI’s Human Users*, 100 *B.U. L. REV.* 1315, 1375 (2020).

²³⁸ Jane Bambauer, *Negligent AI Speech: Some Thoughts About Duty*, 3 *J. FREE SPEECH L.* 343, 346 (2023).

²³⁹ Abbott, *supra* note 92, at 37.

²⁴⁰ See Stein, *supra* note 151, at 982.

²⁴¹ See *id.* at 986.

tion and precaution-taking, and for ensuring that AI users “rise to the level of knowledge about the AI to that of a ‘reasonable person.’”²⁴² In my view Stein’s arguments have merit—and indeed, if products liability in AI cases leans on instructions and warnings as much as I suspect it will,²⁴³ some version of what Stein proposes may in fact begin to develop organically.

Finally, I would be remiss if I failed to mention the several scholars who suspect, as I do, that tort law’s status quo is probably well equipped to handle AI harms. Patrick Hubbard contends that “the legal system’s method of addressing physical injury from robotic machines that interact closely with humans provides an appropriate balance of innovation and liability for personal injury,” and “critiques claims that the system is flawed and needs fundamental change”²⁴⁴ Nicholas Tonti makes a similar argument in the AI-in-healthcare context, and says that “current tort frameworks can accommodate these new dilemmas”²⁴⁵ And over thirty years ago, George Cole suggested that when it comes to harms caused by artificial intelligence, “the best cause of action for most people might be . . . the conduit whereby the reasonableness of human behavior may be evaluated: the cause of action for negligence.”²⁴⁶ I couldn’t agree more.

But whether or not new tort rules would yield more desirable results, there is a more fundamental reason to reject the categorical proposals described in this section. In fact, we have no modern theory of tort law that can support them.

III. THE CASE FOR LETTING AI HARMS LIE

What theory of tort justifies treating AI harms differently than other harms? As the heading suggests, none do—and some AI harms will rightly avoid liability by simple virtue of not being tortious.

Tort law exists to serve a purpose. Reasonable minds (and theories) disagree about what that purpose is—but tort law is unquestionably tethered to *something*. Perhaps the “something” is economic efficiency. Or perhaps it’s righting wrongs, or the fair distribution of losses, or even state-sanctioned vengeance. But whatever the answer, tort law isn’t self-justifying.

²⁴² See *id.* at 1034.

²⁴³ See *supra* Part I.C.

²⁴⁴ F. Patrick Hubbard, “*Sophisticated Robots*”: *Balancing Liability, Regulation, and Innovation*, 66 FLA. L. REV. 1803, 1803 (2014).

²⁴⁵ Nicholas Tonti, *Who to Blame? AI or Your Physician*, 27 QUINNIPIAC HEALTH L.J. 219, 221 (2024).

²⁴⁶ George S. Cole, *Tort Liability for Artificial Intelligence and Expert Systems*, 10 COMPUT./L.J. 127, 213 (1990).

Bearing all this in mind, I begin this Part by proposing two premises that I hope are relatively uncontroversial, and that we can use to analyze the question whether special liability for AI harms is warranted.

First, AI harms should be treated differently than other harms *only if* current theories of tort justify doing so (let's call this premise the *theoretical grounds for special treatment* rule). In other words, if AI harms deserve special treatment, then at least *some* current theory of tort should be able to explain why. If instead the rationale for treating AI harms differently requires a novel theory, then that treatment puts the cart before the horse.

Second, a given tort theory provides a basis for subjecting all AI harms to special treatment *only if* the theory suggests that AI harms *as a category* warrant it (let's call this premise the *categorical* rule). Human harms are not monolithic, and neither are AI harms. We treat human blasters and baristas differently for liability purposes. Eschewing similar distinctions for different AI activities demands a principled basis.

With these two premises in hand—the *theoretical grounds for special treatment* rule and the *categorical* rule—we're now ready to analyze AI harms through the varied lenses of different theories of tort. In the following Sections, I discuss the four most prominent: corrective justice, economic efficiency, civil recourse, and loss distribution.²⁴⁷ But before I do so, a bit of table-setting is in order. Under current tort law, myriad harms are non-redressable—a fact every bit as uncontroversial as it is obvious and socially desirable.²⁴⁸ Let's begin our analysis there.

A. Not All Harms Are Tortious

"The baseline common law regime of tort liability is negligence."²⁴⁹ Negligence has its elements—duty, breach, causation, damages—and where one element is unsatisfied, no tort has occurred.²⁵⁰ Negligence's reasonable person standard stands in the "breach," so-to-speak: whether an actor has breached his duty of care comes down to whether his harm-causing actions were reasonable under the circumstances.²⁵¹

Of course, countless harms are caused by people and companies that acted reasonably, leaving the injured without redress. We know this, and yet in

²⁴⁷ See *infra* Part III.B.–III.E.

²⁴⁸ See *infra* Part III.A.

²⁴⁹ *Ind. Harbor Belt R.R. Co. v. Am. Cyanamid Co.*, 916 F.2d 1174, 1177 (7th Cir. 1990).

²⁵⁰ RESTATEMENT (SECOND) OF TORTS § 281 (A.L.I. 1965). Notwithstanding some other applicable theory of tort liability, of course.

²⁵¹ *Pokora v. Wabash Ry. Co.*, 292 U.S. 98, 104 (1934) ("Standards of prudent conduct are declared at times by courts, but they are taken over from the facts of life.").

general we have not risen up against it.²⁵² Far from it. With vanishingly few exceptions, in the United States we assume the risk of faultless injury by living in and navigating a world full of casual danger.

In other words, fault is the enemy we've chosen to fight; injury is merely an element. What to make, then, of AI actors that can mindlessly do things that we've grown accustomed to weighing the blameworthiness of? Perhaps this is why AI harms are so unnerving to us—and why many have reflexively reached for the strict-liability elephant gun that we usually keep locked away in the attic.

Lack of remedy is a frequent feature of our system. If that feature is present for many AI harms, too, it doesn't automatically turn into a bug. Some injurers are blameless; some risks are assumed; some harms are socially acceptable.

Keeping all this in mind, let's consider AI harms in the context of the corrective-justice theory of tort law.

B. Corrective Justice Theory

Corrective justice is principally concerned with wrong-righting.²⁵³ As Ernest Weinrib notes, corrective justice is focused on the “normative relationship between the parties.”²⁵⁴ Under a corrective-justice theory of tort, “liability reflects the conclusion that the defendant and the plaintiff have respectively done and suffered the same injustice.”²⁵⁵ Where there is liability, there is a moral wrong—and corrective justice “asserts a connection between the remedy and the wrong.”²⁵⁶

All this looks familiar, of course—and boils down to fault. Indeed, corrective justice theorists view fault as “central” to tort law and its “moral defensibility.”²⁵⁷ Tort law in practice—and for the most part, how we teach it—is consistent with a corrective justice approach. Wrongs resulting in losses or unjust gains are “rectified, eliminated, or annulled”—and what constitutes a redressable wrong is defined not only by the harm, but by the characteristics of, and

²⁵² In fact, some have, and compellingly. But economically sound though these positions might be, fifty years on, they still haven't revolutionized tort law. *See, e.g.*, Epstein, *supra* note 156, at 188 (arguing strict liability for virtually all harms is preferable to negligence, because strict liability would still be economically efficient, and doing so would greatly “reduce the administrative costs of decision”).

²⁵³ ERNEST J. WEINRIB, *CORRECTIVE JUSTICE* 16 (Timothy Endicott, John Gardner & Leslie Green eds., 2012) (“As its name indicates, corrective justice has a rectificatory function. By correcting the injustice that the defendant has inflicted on the plaintiff, corrective justice asserts a connection between the remedy and the wrong.”).

²⁵⁴ *Id.* at 9.

²⁵⁵ *Id.* at 10.

²⁵⁶ *Id.* at 16.

²⁵⁷ Jules L. Coleman, *Tort Law and the Demands of Corrective Justice*, 67 *IND. L.J.* 349, 351 (1992).

context surrounding, the harm-causer's actions.²⁵⁸ Simply put, "corrective justice is concerned with the gains and losses *that one person causes another*."²⁵⁹

This is enough to permit the conclusion that corrective justice offers no principled grounds for imposing special liability on AI harms. But we can say more: corrective justice implies that many AI harms shouldn't be subject to tort liability at all.²⁶⁰ It is not enough to say that AI actors might cause injuries that "need" to be redressed, or that the existence of amoral machines is itself a kind of evil. To a corrective justice theorist, "[e]vil and need are moral categories that may well figure in other normative contexts, but they are not pertinent to liability."²⁶¹ Corrective justice is about righting wrongs between people—and neither virtue nor need are sufficient to "connect any two particular persons" to a moral wrong in need of redress.²⁶² Accordingly, an amoral AI actor causing harm would not, alone, require redress under a corrective justice theory. Corrective justice would invite recourse only if it could be shown that a person deployed the harm-causing AI in a blameworthy, wrong-causing way.²⁶³

So corrective justice doesn't support special liability for AI harms. What's more, it suggests that AI harms not directly traceable to human fault are unreachable. (So much for strict liability.) How about economic efficiency theory?

C. Economic Efficiency Theory

Starting in the early 1960s, scholars began to develop an economic theory of tort law. Ronald Coase broke ground when he introduced what we now call the "Coase Theorem,"²⁶⁴ and soon after, Guido Calabresi published his first of several writings on the economics of tort law.²⁶⁵ A decade after that, Calabresi's book *The Cost of Accidents*²⁶⁶ incited Richard Posner—among law and economics' best tort-law thinkers²⁶⁷—to join the fray. He and William Landes

²⁵⁸ *Id.* at 357.

²⁵⁹ *Id.* at 358 (emphasis added).

²⁶⁰ *Id.* (explaining that corrective justice is concerned with redressing harms one individual causes another); see also WEINRIB, *supra* note 253, at 16 (noting that corrective justice remedies "a connection between the remedy and the wrong").

²⁶¹ WEINRIB, *supra* note 253, at 19.

²⁶² *Id.*

²⁶³ See *id.*

²⁶⁴ See generally R.H. Coase, *The Problem of Social Cost*, 3 J.L. & ECON. 1 (1960) (articulating the economic theory now known as the Coase Theorem).

²⁶⁵ See generally Guido Calabresi, *Some Thoughts on Risk Distribution and the Law of Torts*, 70 YALE L.J. 499 (1961) (considering the economics of resource allocation in determining whether the basis of liability should be fault or risk distribution).

²⁶⁶ See generally Richard A. Posner, Book Review, 37 U. CHI. L. REV. 636, 636–48 (1970) (reviewing GUIDO CALABRESI, *THE COST OF ACCIDENTS: A LEGAL AND ECONOMIC ANALYSIS* (1970)); RICHARD A. POSNER, *ECONOMIC ANALYSIS OF LAW* (1972).

²⁶⁷ See, e.g., Saul Levmore, *Richard Posner, the Decline of the Common Law, and the Negligence Principle*, 86 U. CHI. L. REV. 1137, 1137 (2019) (discussing Posner's influence on law and economics in tort law).

later published a book-length work setting forth a positive economic theory of tort law.²⁶⁸ Despite its positive-law focus, that work in particular offers guidance on how AI harms ought to be treated.

From a law and economics perspective, tort law is best understood by analyzing the interplay between utility maximization and economic efficiency.²⁶⁹ Actors are presumed to be rational, and therefore to seek, in all situations, to maximize their own utility.²⁷⁰ Utility encompasses any number of things a person might rationally care about: welfare, happiness, satisfaction, etc.²⁷¹ Because utility's contents are many and hard to measure, we treat utility as dependent on one composite good: income.²⁷² More income is always better. So gains in utility are represented by gains in income, whereas losses in utility—harms—are represented by losses in income.²⁷³ Because individuals act to maximize their own utility, we can conduct a kind of game-theory analysis²⁷⁴ to determine what rules are most economically efficient. The efficient rule is the one that maximizes *total* utility (add everyone's utility up).²⁷⁵

We have a bit more table-setting to do, but already we can see that economic efficiency has a big leg-up over corrective justice when it comes to analyzing AI harms. Because economic efficiency weighs utility rather than rectitude, it can account for *amoral* actions as readily as immoral ones. This has important implications. For one thing, from a utility-maximization point of view, the best ex-ante liability regime might plausibly be to ignore fault altogether and make strict liability the universal rule.²⁷⁶

As it turns out, most law and economics theorists agree that from an efficiency standpoint, there isn't much difference between negligence and strict liability. Under either liability regime, actors will adopt the same due-care level.²⁷⁷ As Richard Epstein explains, the theories diverge “only with regard to

²⁶⁸ See LANDES & POSNER, *supra* note 165, at 9 (stressing that the authors are focused on a positive economic theory of torts, and therefore are “interested in explaining, rather than defending, the common law of torts”).

²⁶⁹ *Id.* at 1, 19.

²⁷⁰ *Id.* at 32.

²⁷¹ *Id.* at 54.

²⁷² *Id.*

²⁷³ *Id.* at 54–55.

²⁷⁴ See Randal C. Picker, *An Introduction to Game Theory and the Law* 11–13 (Coase-Sandor Inst. for L. & Econ., Working Paper No. 22, 1994) (applying game theory to tort liability theories).

²⁷⁵ LANDES & POSNER, *supra* note 165, at 16 (“A change is wealth maximizing if the dollar value of the gains to the winners is greater than the dollar cost of the losses to the losers. The positive economic theory of tort law holds that tort rules are efficient in the sense of wealth maximizing.”).

²⁷⁶ See Epstein, *supra* note 156, at 188 (highlighting that strict liability rules reduce administrative costs).

²⁷⁷ See LANDES & POSNER, *supra* note 165, at 64 (“Many legal analysts think that strict liability would induce potential injurers to be more careful Economic analysis suggests that this view is superficial. Assuming that $x^* = 0$, B will choose y^* , the due care level, whether he is strictly liable or liable only if negligent.” (footnote omitted)).

that portion of the damages against which no prudent precautions could have been taken.”²⁷⁸ What’s more, the transaction costs of litigating a given strict-liability claim are lower than litigating a claim under negligence.²⁷⁹

Of course, this isn’t quite the whole story. For one thing, strict liability will likely result in many more claims than negligence.²⁸⁰ Determining which regime has higher transaction costs overall will require a comparison of claim frequency and costs under each standard.²⁸¹ There’s also the uncertain question of victim activity levels. Posner and Landes point out that predicting the effects of a given regime on victim activity is both difficult and highly relevant to ultimate efficiency.²⁸² (Highly relevant because more victim activity means more claims—see above.) For this and other reasons they view it as impossible to pronounce one regime across-the-board superior to the other.²⁸³

One more thing: even supposing strict liability and negligence are equally efficient, what they incentivize differs in a key respect. Although both negligence and strict liability provide an incentive to exercise due care, strict liability additionally incentivizes *reduced activity* on the part of the potential defendant, and might well incentivize potential plaintiffs to act more—and more carelessly.²⁸⁴ This provides an economic-efficiency account for why we deviate from negligence when it comes to abnormally dangerous activities.²⁸⁵ For these, we want to motivate potential defendants to achieve their goals differently—not merely with greater care.²⁸⁶

Now we have enough to work with. From here, let’s pull in as priors some of our discussion from Part I.²⁸⁷ First, AI actors cannot themselves be negligent.²⁸⁸ Second, under products-liability theory, design and manufacturing defects might be hard to find, let alone prove.²⁸⁹ To these let’s add our *categori-*

²⁷⁸ Epstein, *supra* note 156, at 188.

²⁷⁹ See *id.* (noting strict liability rules can reduce administrative costs). For example, under strict liability, parties needn’t litigate the often difficult “due care” issue.

²⁸⁰ See LANDES & POSNER, *supra* note 165, at 65 (“Under strict liability a claim arises every time there is an accident caused by someone worth suing; under negligence there is a claim only if the victim thinks that he can show that the defendant failed to use due care.”).

²⁸¹ See *id.* at 65–66 (noting the potential differences in numbers of claims and information costs under negligence and strict liability frameworks).

²⁸² *Id.* at 69–71.

²⁸³ *Id.* at 71.

²⁸⁴ *Id.* at 66–69 (highlighting that strict liability may provide an incentive to reduce activity).

²⁸⁵ *Id.* at 70 (noting that if seeking a change in defendant’s activity level is an efficient way to avoid accidents, strict liability should be the rule chosen).

²⁸⁶ *Id.* at 66 (discussing a key difference in negligence and strict liability being an incentive to reduce accidents via change in level activity rather than increasing care with which one performs the activity).

²⁸⁷ See *supra* Part I.

²⁸⁸ See *supra* Part I.B.

²⁸⁹ See *supra* Part I.C.

cal rule from the top of this Part: categorically subjecting AI harms to special treatment is justified *only if* AI harms as a category warrant it.²⁹⁰

Under an economic efficiency theory, AI harms would warrant a special liability rule if and only if implementing one would increase total utility over the current rule.²⁹¹ But even if this were true in cherry-picked cases, it wouldn't be true across the board—violating our *categorical* rule.

Let's look at a standard case first. Take the makers of household goods implementing AI features—recall our self-sharpening knife, stability-detecting ladder, and speed-adjusting treadmill from Part II.A.2.²⁹² As we already discussed, those devices showed no evidence of defectiveness; they were designed reasonably and manufactured properly according to their design.²⁹³ Nor would placing them in the stream of commerce be negligent, assuming no outlandish marketing claims. (“With this ladder, climb confidently *without* a safety harness!”—nonsense like that.) What effect would a more stringent liability regime have on the total utility of products like these?

A negative one. These products could hardly be made or sold more safely, and recall Epstein's callout that negligence and strict liability diverge “with regard to that portion of the damages against which no prudent precautions could have been taken.”²⁹⁴ So a more stringent liability regime would reduce the user's incentive to be careful and force the seller to do one of three things: eat the increased cost of liability; raise prices (eating only some of the costs, depending on the price elasticity of demand²⁹⁵); or stop selling. Which the seller chooses will be a function of the market. If demand is relatively inelastic, the seller will raise prices.²⁹⁶ But if raising prices would tank demand (high elasticity), the seller will raise prices only a little, or not at all. And if the costs of liability would exceed profits regardless of price, the seller will simply pull the product off the market.²⁹⁷

If we accept that these products were made and sold as safely as is socially desirable, then these results are inefficient. Knives—even AI self-sharpening ones—can be dangerous when carelessly used. Same with AI ladders and treadmills. Here, the users of these safely-made-and-sold products are the ob-

²⁹⁰ See LANDES & POSNER, *supra* note 165, at 66 (“Under strict liability a potential injurer will interpret y^* broadly. He is interested in any measure that would reduce his expected damage by more than the cost of the measure; it is a matter of indifference to him whether the cost results from purchasing some safety input or from foregoing the profits from a higher level of productive activity.”).

²⁹¹ See *id.* at 16 (discussing the positive economic theory of tort law as equating efficient tort rules with wealth maximization).

²⁹² See *supra* Part II.A.2.

²⁹³ See *supra* Part II.A.2.

²⁹⁴ Epstein, *supra* note 156, at 188.

²⁹⁵ See Hylton, *supra* note 185, at 2480 n.94 (“In general, the extent to which the costs are passed on the short run is a function of demand and supply elasticities.”).

²⁹⁶ *Id.* at 2479–80.

²⁹⁷ *Id.*

vious least-cost avoiders.²⁹⁸ And just as is true with any other safely-made-and-sold product, economic efficiency suggests we ought to choose the liability regime that incentivizes efficient accident avoidance.²⁹⁹ As it happens, negligence and products liability, coupled with comparative negligence, accomplishes this—and is the standard framework almost everywhere.³⁰⁰

Now for a few cherry-picked cases that engage with the worries this article seeks to respond to. First, let's put AI drivers in the hotseat once more. Recall from Part I that negligence for driverless cars *is* still viable: even if AI actors aren't susceptible to negligence suits, their legal-entity deployers are.³⁰¹ As we discussed in Part I, although it may not be clear today just how safe a driverless vehicle must be to satisfy a deployer's duty of care, courts are perfectly capable of answering that question.³⁰² What's more, driverless vehicles have many characteristics that potentially increase total utility³⁰³—providing strong reasons to avoid disincentivizing their existence and development. All this suggests that negligence—and not some more stringent regime like strict liability—remain appropriate from an economic-efficiency standpoint.

But consider a higher-stakes example: an AI actor designed to manage core functions at a nuclear powerplant just outside a major U.S. city.³⁰⁴ Let's stipulate the obvious—that if the AI screws up, the consequences will be incalculably bad.

No need to bury the lede on this one. From an economic-efficiency point of view, strict liability, not negligence, is the proper standard.³⁰⁵ First, the potential harm to victims is astronomical—and apart from moving or building fallout shelters, there isn't much a potential plaintiff could do to mitigate it. Second, as a result of the degree of harm, the duty of care would effectively be perfection—so negligence would always reach the same result as strict liability.

²⁹⁸ See STEVEN SHAVELL, *ECONOMIC ANALYSIS OF ACCIDENT LAW* 17 (1987) (“The notion of the least-cost avoider applies in situations where the risk of accidents will be eliminated if *either* injurers or victims take care. In such situations it is clearly wasteful for *both* injurers and victims to take care; rather, it is optimal for the parties who can prevent accidents at least cost, the ‘least-cost avoiders,’ alone to take care.”).

²⁹⁹ LANDES & POSNER, *supra* note 165, at 62–64.

³⁰⁰ See *supra* Part I.B.

³⁰¹ See *supra* Part I.B.

³⁰² See *supra* Part I.B.

³⁰³ See Ed Garsten, *What Are Self-Driving Cars? The Technology Explained*, FORBES, <https://www.forbes.com/sites/technology/article/self-driving-cars/> [perma.cc/EV6R-LHMR] (Jan. 24, 2024) (highlighting benefits of self-driving cars, including improved safety, improved ride comfort, and the ability to conduct other utility-generating activities while driving).

³⁰⁴ No doubt this would be governed by federal regulation—but for discussion's sake, let's pretend the common law has to work out the proper rule. See 42 U.S.C. § 2014(w) (providing, at the federal level, the definition of “public liability,” which includes any liability from a nuclear incident).

³⁰⁵ LANDES & POSNER, *supra* note 165, at 63.

ity,³⁰⁶ just with higher information costs.³⁰⁷ This is precisely the sort of situation Posner and Landes note is an appropriate fit for strict liability,³⁰⁸ and it fits nicely into any standard analysis of abnormally dangerous activities, to boot.³⁰⁹

Not surprisingly, real-world harms that result from nuclear power plants are effectively subject to strict liability, imposed by statute.³¹⁰ For a more workaday example that involves the same considerations, consider urban demolitions. AI robot blasters will probably make the demolition and construction industries safer. But even if robots are doing all the work, blasting will still be abnormally dangerous—and still warrant strict liability. What does this mean? Merely the obvious: that, just like some human activities, some AI activities will warrant strict liability, too.

The upshot of this whole subpart is that categorical special liability for all AI harms is unsupported by economic efficiency theory. Next at bat is civil recourse theory—which, as we’ll see, gets us no further than, and has the same AI-actor blind spot as, corrective justice.

D. Civil Recourse Theory

Civil recourse theory has much in common with corrective justice.³¹¹ Both are relational³¹²; both handle wrongs³¹³; both reject the economic account of tort law.³¹⁴ But civil recourse theorists would say that it differs from corrective justice in a key respect. Under corrective justice theory, tort law rights

³⁰⁶ Meaning there is no liability gap. See Epstein, *supra* note 156, at 188 (explaining that “strict liability diverge[s] from . . . negligence only with regard to that portion of the damages against which no prudent precautions could have been taken”).

³⁰⁷ LANDES & POSNER, *supra* note 165, at 66–67.

³⁰⁸ *Id.* at 63.

³⁰⁹ See, e.g., RESTATEMENT (SECOND) OF TORTS § 520 (A.L.I. 1977) (providing the factors in determining whether an activity is abnormally dangerous).

³¹⁰ See 42 U.S.C. § 2014(w) (defining “public liability” to include any legal liability from a nuclear incident).

³¹¹ Indeed, there are convincing arguments that civil recourse is, at most, a subset of corrective justice. But for completeness I give them separate treatment here. See generally Scott Herschovitz, *Corrective Justice for Civil Recourse Theorists*, 39 FLA. ST. U. L. REV. 107, 109 (2011) (“[C]ivil recourse is best understood as a corrective justice account of tort.”); Ernest J. Weinrib, *Civil Recourse and Corrective Justice*, 39 FLA. ST. U. L. REV. 273, 273 (2011) (“The theories of civil recourse and corrective justice are so closely related that when Ben Zipursky was in Toronto several years ago presenting his paper *Civil Recourse, Not Corrective Justice*, I publicly asked him whether the word ‘not’ in the title was a typo.”).

³¹² See Benjamin C. Zipursky, *Civil Recourse, Not Corrective Justice*, 91 GEO. L.J. 695, 701 (2003) (“In the sense that the defendant and plaintiff are two poles of a relationship within which the compensation travels, and within which it stays, tort law is intrinsically bipolar.”).

³¹³ *Id.* at 697–98 (“[T]he model of ‘rights, wrongs, and recourse’ . . . differs strikingly from corrective justice theory in several key respects. Most prominently, this view does not utilize the notions of rectification, normative equilibrium, or a duty of repair.”).

³¹⁴ *Id.* at 698.

wrongs by assigning the defendant a duty to “make the injured party whole.”³¹⁵ But under civil recourse theory, tort law works by granting plaintiffs a state-sanctioned right “to act against the defendant.”³¹⁶

To a civil recourse theorist, “what the defendant has done [is] ultimately subsidiary to questions about what the plaintiff is entitled to get.”³¹⁷ Superficially it might seem like focusing on the rights of the injured party would make civil recourse more capable of directly handling AI harms. After all, corrective justice had trouble handling harms caused by AI actors in part because AI actors can’t be held morally to blame.³¹⁸ But in fact, civil recourse cares at least as much about blameworthiness as corrective justice does: any right of action by the plaintiff is predicated on their “[h]aving been wronged by the defendant.”³¹⁹

When Benjamin Zipursky first proposed civil recourse theory, he anticipated and rejected the criticism that civil recourse recasts tort law as state-sanctioned vengeance.³²⁰ But he nonetheless proposed that tort law does permit “getting even,” and that the ability to do so might even be “an important function of our tort system.”³²¹ Whether the state-sanctioned-vengeance critique is fair or not, it helps illuminate why civil recourse theory does not permit AI actors to be dealt with directly by tort law, and therefore provides no new basis for treating AI actors or their harms differently. A state-sanctioned right of action against an AI actor, whether for vengeance or some other purpose, would leave plaintiffs as empty-handed as would kicking a bedpost.³²² Of course, under civil recourse theory, a deployer who wrongfully deploys an AI remains reachable through tort law; civil recourse has no relevant bones to pick with our pre-existing negligence and products liability regimes.³²³ That being said, strict liability sits just as uneasily with civil recourse as it does with corrective justice—and civil recourse certainly provides no grounds for imposing it or some other special liability rule for AI harms.

What’s left is the loss distribution theory of tort. Is it the last best hope for a theoretical justification for special liability for AI harms? Let’s find out.

³¹⁵ *Id.* at 695.

³¹⁶ *Id.* at 734; *see id.* at 699 (“In permitting and empowering plaintiffs to act against those who have wronged them, the state is not relying upon the idea that a defendant has a pre-existing duty of repair. Instead, it is relying on the principle that plaintiffs who have been wronged are entitled to some avenue of civil recourse against the tortfeasor who wronged them.”).

³¹⁷ *Id.* at 733.

³¹⁸ *See supra* Part III.B.

³¹⁹ Zipursky, *supra* note 312, at 735; *see id.* at 737 (“I am articulating the concept of a right of action based on having been wronged and examining what commitments are entailed by recognition of this concept in our system.”).

³²⁰ *Id.* at 737.

³²¹ *Id.*

³²² *See id.* at 733 (stating that “what the defendant has done [is] ultimately subsidiary to questions about what the plaintiff is entitled to get”).

³²³ *See id.* at 735, 737 (noting that a cause of action is still “based on having been wronged”).

E. Loss Distribution Theory

Unlike corrective justice or civil recourse, loss distribution is little concerned with fault.³²⁴ And unlike economic efficiency, it is unconcerned with total utility, let alone maximizing it.³²⁵ (Indeed, loss distribution is vulnerable to the criticism that it does little more than propose to turn tort law into a highly inefficient insurance program.) Instead, loss distribution is concerned with one thing: ensuring victim compensation.³²⁶

Loss distribution theory has been out of vogue for quite some time now, but its influence persists in discrete pockets of tort law. Workers' compensation is one example;³²⁷ the strict-liability standard we hold manufacturing defects to is arguably another.³²⁸ In the view of its early-twentieth-century proponents, the principal function of tort law was, simply, to compensate the injured.³²⁹ To them, tort law was little more than a kind of "public law in disguise";³³⁰ basically, a social insurance regime.³³¹

³²⁴ Fleming James Jr., *Contribution Among Joint Tortfeasors: A Pragmatic Criticism*, 54 HARV. L. REV. 1156, 1158–59 (1941).

³²⁵ See George L. Priest, *The Invention of Enterprise Liability: A Critical History of the Intellectual Foundations of Modern Tort Law*, 14 J. LEGAL STUD. 461, 470 (1985) ("The single idea on which all of James's tort scholarship was built was the role of personal injury damage judgments as a form of 'social insurance.' . . . James was greatly concerned about the consequences of serious personal injury on the lives of victims, most particularly the poor." (footnote omitted)); LANDES & POSNER, *supra* note 165, at 57–58 ("If people who want insurance and are willing to pay for it can obtain it in the insurance market or in some informal substitute, there is no apparent reason to use the tort system to provide insurance also. . . . We can put this a little more strongly: the tort system is an exceedingly costly insurance mechanism." (footnote omitted)).

³²⁶ See James Jr., *supra* note 324, at 1158 (questioning fault in accident cases, and describing the fault principle as "overshadow[ed]" by the need to ensure compensation, and the need for "expedient allocation of their losses so that society may absorb them with as little harm to itself as possible").

³²⁷ See Priest, *supra* note 325, at 472–76 (discussing James's hope for a comprehensive compensation plan).

³²⁸ See *id.* at 502 (discussing James's compensation system in relation to liability for manufacturing defects).

³²⁹ See *id.* at 470–71 ("James . . . conceived the principal function of tort law to be . . . compensation of the injured.").

³³⁰ Leon Green, *Tort Law Public Law in Disguise: II*, 38 TEX. L. REV. 257, 269 (1960) ("[A]ny judgment arrived at is nothing more than the reactions of the men called upon to give judgment in the particular case as they conceive to be desirable for all parties concerned—the parties litigant, other parties who may have similar cases, and all the rest of us who may have a stake in what we call a just decision. In this sense tort law calls for as delicate and as profound probing as any of the other areas of law in which the public good is more obvious. It is thus that I conclude that tort law is public law in disguise.").

³³¹ James Jr., *supra* note 324, at 1157 ("The full blessings of distribution can best be attained by comprehensive social insurance, but some measure of advantage may be had within the present framework of our law of torts, for many of its rules tend in practice to shift losses to agencies which can and do distribute them."); see also LANDES & POSNER, *supra* note 165, at 5 ("Legal realists such as Fleming James . . . saw the only attainable function of tort law as the provision of social insurance.").

Loss distribution theorists no doubt believe that AI harms ought to be redressed through tort law irrespective of fault, under something like strict liability, or some formulation of enterprise or common-enterprise liability. (This is roughly what some of the scholars discussed in Part I have proposed.)³³² But loss distribution theory proposes similar treatment for *all* accident-related injuries; it provides no basis for treating harms caused by AI actors differently than other kinds of harm.³³³ In other words, loss distribution fails to satisfy our *theoretical grounds for special treatment* rule. To a loss-distribution theorist, injuries caused by AI harms are not theoretically distinct from other kinds of injuries—and so do not justify special treatment.

Now for a partial concession. Though I am unconvinced loss distribution theory provides grounds for treating AI harms differently than other harms, I am at least amenable to the argument that AI harms offer novel support for loss distribution theory. We may well be entering a technological epoch in which fault will become somewhat irrelevant. In the future, perhaps we won't be scratching our collective heads on AI liability, but rather, remembering the current moment as the age of accidents. If our future is packed full of super-safe computers, costly accidents may all but vanish—and with them, most tort suits that deal in physical injury. Ryan Abbott makes precisely this point: "We could soon be living in a world where no one dies from unintended injury, or from medical error for that matter."³³⁴ One might imagine that in this potential future, a full rethink of tort law might, on practical grounds, be worthwhile.

Then again, the future I just described seems ripe for a different shift altogether—one that moves away from tort law, and toward private insurance (the availability of which, after all, is the best rejoinder to loss distribution).³³⁵ That is the primary topic of the next and final Part.³³⁶

IV. INSURANCE, RISK ASSUMPTION, AND AI CATASTROPHES

This Part first discusses insurance as a solution to addressing the uncertainty of liability that comes with AI.³³⁷ It then moves on to argue that, even

³³² See Vladeck, *supra* note 15, at 129 n.39 ("[A]s I envision it, common enterprise liability here would be a form of court-compelled insurance. Manufacturers and designers . . . would jointly indemnify individuals insured by driver-less cars when it is impossible to determine fault.").

³³³ See James Jr., *supra* note 324, at 1158 (highlighting the primary concern of loss distribution theory being victim compensation, regardless of the cause or type of harm).

³³⁴ Abbott, *supra* note 92, at 44.

³³⁵ See SHAVELL, *supra* note 298, at 208 ("[W]here victims can secure accident insurance coverage (or can be supplied with social insurance), the main advantage of the liability system is that it provides injurers incentives to reduce risk."); see also LANDES & POSNER, *supra* note 165, at 178.

³³⁶ See *infra* Part IV.

³³⁷ See *infra* Part IV.A.

for AI catastrophes, for which insurance may be insufficient, negligence stands ready to handle them.³³⁸

A. Insurance Markets, Old and New

While the tort system offers effective mechanisms for setting incentives and redressing moral wrongs, as a form of social insurance it leaves much to be desired, thanks to its sky-high transaction costs.³³⁹ Consider just a few of them. First, suing stinks. The emotional costs are often extremely high and difficult to measure.³⁴⁰ And waiting stinks, too: your average personal-injury lawsuit takes months or years to reach a settlement or verdict, with some states' civil dockets averaging around five years per case.³⁴¹ Now suppose you win or settle. What will you actually get? Roughly sixty percent of the award, maybe less; your lawyer gets the rest.³⁴² (And that's assuming you sued directly. If you're a part of a class action, your meager payout might not even be worth the hassle of cashing the check.)

Again, there are many good reasons for having a tort system like we have—even one that sometimes requires us to go through the above rigamarole. (Recall our earlier theory discussions about economic efficiency and civil recourse.)³⁴³ But for injuries that don't trigger tort law—for example, accidents that occur in the absence of negligence or defect—insurance provides a much more efficient vehicle for compensation.³⁴⁴

For many kinds of harms caused by AI actors, we likely don't even need *new* insurance markets. For example, the health insurance market is just as capable of compensating for health-related harms caused by AI as those arising in any other way. True, insurers' actuaries will likely struggle initially to cali-

³³⁸ See *infra* Part IV.B.

³³⁹ See LANDES & POSNER, *supra* note 165, at 57–58 (highlighting that the tort system is costly).

³⁴⁰ See Peter Bowal & Jacqueline Bowal, *Small Claims Court: A Venue Made for Self-Represented Litigants*, LAWNOW (June 30, 2015), [https://www.lawnow.org/small-claims-court-a-venue-made-for-self-represented-litigants/\[perma.cc/L4NL-WQNX\]](https://www.lawnow.org/small-claims-court-a-venue-made-for-self-represented-litigants/[perma.cc/L4NL-WQNX]) (“There are . . . non-economic costs of suing such as adverse publicity, stress and mental anguish, uncertainty, time to prepare and prosecute the case, and destruction of relationships. We think you should not sue unless it is for a relatively large amount of money and you have a good chance of getting paid if successful.”).

³⁴¹ See Michael Heise, *Justice Delayed?: An Empirical Analysis of Civil Case Disposition Time*, 50 CASE W. RES. L. REV. 813, 836–38 (2000) (providing data on case disposition time in various counties).

³⁴² See Lester Brickman, *Contingent Fees Without Contingencies: Hamlet Without the Prince of Denmark?*, 37 UCLA L. REV. 29, 30 (1989) (“According to conventional wisdom virtually all contingent fee percentages exceeding fifty percent are illegal and excessive, but most lower percentages are valid.”); see, e.g., *In re Compact Disc Minimum Advertised Price Antitrust Litig.*, 370 F. Supp. 2d 320, 321 (D. Me. 2005) (noting that individual payouts were roughly \$4, and only two percent of the class made a claim).

³⁴³ See *supra* Part III.B–C.

³⁴⁴ See, e.g., Abbott, *supra* note 92, at 30–31 (discussing the potential use of insurance to cover AI liability).

brate around this new source of potential liability. But because AI actors have the potential to increase safety and reduce the number and severity of accidents overall,³⁴⁵ thanks to AI, health insurers are likely to experience a net decrease in costs.³⁴⁶

So when it comes to medical-related AI injuries, health insurance can probably cover those just fine. What about other kinds of harms? Let's say an AI investment tool causes my portfolio to underperform the market by twenty percent. Or perhaps a fully autonomous bicycle crashes into and destroys my mailbox. Or maybe an AI housekeeping robot throws away one of my family heirlooms. Assume neither negligence nor products liability can reach these harms (for any of the reasons we discussed in Part I).³⁴⁷ What to do?

Some harms—like the one caused by the AI investment tool above—are neither practicably insurable nor deserving of redress. Remember, this tool is not defective, so it no doubt carries sufficient warnings about what the tool can and cannot do, and about the inherent risks of investing. In this situation, by using the AI investor to engage in an inherently risky activity (stock market investing), I've simply assumed the risk.³⁴⁸

But what about those other harms—my mailbox, destroyed by the autonomous bicycle; or the robot housekeeper that chucked my heirloom? Insurance could handle things like these. Perhaps homeowners' insurance companies begin offering an AI rider.³⁴⁹ There's no reason State Farm can't cover my mailbox—and AI harms to property will probably be rare enough that I could obtain the rider pretty cheaply.³⁵⁰ As for the heirloom, that's tricky in the same way that insuring anything with sentimental value is.³⁵¹ If the heirloom is a

³⁴⁵ See *id.* at 44 (stating technological advancements could reduce harm in such a way that renders the discussion surrounding negligence and strict liability unnecessary).

³⁴⁶ For other reasons too, including the administrative efficiencies that AI is already beginning to provide. See Waqaas Al-siddiq, *Accelerating Healthcare with AI: Reducing Administrative Burdens*, FORBES (Jan. 7, 2025), <https://www.forbes.com/councils/forbesbusinesscouncil/2025/01/07/accelerating-healthcare-with-ai-reducing-administrative-burdens/> [perma.cc/8665-B5SU] (“AI holds immense promise in alleviating this stress on the healthcare system. By streamlining and automating administrative tasks, AI can reduce human error, improve operational efficiency and allow healthcare providers to focus on direct patient care.”).

³⁴⁷ See *supra* Part I.

³⁴⁸ See Stein, *supra* note 151, at 991 (“Primary implied assumption of risk applies when the risk encountered is inherent in the activity itself; it is ‘neither created nor exacerbated by negligence’ but ‘simply exists.’” (citation omitted)).

³⁴⁹ *What Is an Insurance Rider on a Homeowners, Condo, or Renters Policy?*, PROGRESSIVE, <https://www.progressive.com/answers/insurance-rider/> [perma.cc/2G92-7WZW] (“A[n] [insurance] rider is an add-on to a homeowners, renters, or condo insurance policy.”).

³⁵⁰ Though admittedly, it's doubtful my deductible would make a claim worthwhile. See *What Is Home Insurance Coverage?*, STATEFARM, <https://www.statefarm.com/insurance/homeowners/home-insurance-coverage> [perma.cc/7C6K-A6M2] (providing information about StateFarm's home insurance coverage).

³⁵¹ See *id.* (discussing home insurance coverage as helping to cover costs of heirlooms).

piece of jewelry that I value only as an investment, insurance works well. (And jewelry and valuable-items insurance exists already! Once again: AI rider.) But if the heirloom is a photo album or baby blanket, insurance won't get me very far—so I'm probably out of luck. This is sad but probably doesn't justify compensation. (Though a one-star product review might be in order, depending on how well my robot housekeeper handles my laundry and dishes.)

In short, current forms of insurance are already well positioned to (profitably!) insure against AI harms. But this doesn't mean we won't see new forms of AI-specific insurance crop up. Indeed, given the delta between the likely low actual risk of AI harms and the general population's high perceived risk of them, at some point soon I would expect companies to begin offering AI-specific insurance. Perhaps it will take the form of umbrella insurance, designed to cover what standard insurance plans—say, health, home, and car insurance—don't.³⁵² Or maybe we'll see AI insurance tailored to compensate for harms caused only by others' use of AI. Or perhaps even AI-consumer-products insurance, designed to give us assurance that AI robots really can be trusted not to break our dishes, or vacuum up the family cat. Only time will tell.

B. AI Catastrophes: Reprising Negligence

Negligence is a timeless doctrine. Its chief innovation, the reasonable-person standard, is flexible. And when paired with concepts like the Hand formula³⁵³ and least-cost avoidance,³⁵⁴ it is powerful. Negligence can handle far more than the careless deployment of an obviously broken AI. It is equally capable of handling the unreasonable deployment of an AI that, though technologically flawless, is too powerful to be set loose. Reasonableness is a function of decisions and context, and the advent of incomprehensibly potent artificial intelligence is rapidly changing that context. Negligence can handle it—and so, too, can the courts, judges, and juries that will be responsible for applying the law to the facts. We ought to trust the systems we've built, and the theoretical foundations we've built them on.

Speaking of our systems, tort law is only one of the ways that we prevent and redress risk and harm. Where artificial intelligence enters spaces that pose unique risks—critical infrastructure, for example, or airline piloting—legislators and agencies are well equipped to manage this through targeted, activity- and context-specific legislation and regulation. This is how we manage human ac-

³⁵² See *Umbrella Insurance—How It Works & What It Covers*, GEICO, <https://www.geico.com/information/aboutinsurance/umbrella/> [perma.cc/SUH6-72ZN] (defining umbrella insurance as “provid[ing] coverage beyond the limits of your other insurance policies”).

³⁵³ See LANDES & POSNER, *supra* note 165, at 85–86 (explaining the Hand Formula).

³⁵⁴ See SHAVELL, *supra* note 298, at 17 (explaining the idea of least-cost avoider).

tivities, after all, and there is every reason to believe this approach will be similarly effective for managing AI activities as well.

CONCLUSION

More than a century ago, New Yorkers panicked over an invisible electrical current they could neither see nor understand. Yet the law kept pace—not by re-writing liability rules from the ground up, but by applying the ordinary duty of reasonable care to an extraordinary new technology. The sky never fell, and the lights stayed on. AI provokes a similar species of wonder and dread. And it tempts us toward the same blunt prescriptions—strict liability, enterprise liability, no-fault funds—that once again propose to tame the unfamiliar by discarding fault. This article has tried to show that those promises, and the problem they propose to solve, are mostly illusory.

Products liability already supplies the tools needed for the subset of AI harms traceable to defective design, manufacturing, or warning. Negligence, fortified by familiar evidentiary shortcuts and burden-shifting devices, can reach the rest. And none of the established justificatory frameworks—corrective justice, economic efficiency, civil recourse, or loss-distribution—furnish a principled reason to brand AI harms as *sui generis*. Where plaintiffs can prove wrongful deployment or defectiveness, there will be liability. Where they can't, there won't be—and better tools than tort will manage the losses.

Tort, after all, is not our only safety net. Competitive insurance markets have historically leapt into every gap opened by new technologies, and early signals suggest they will do the same for AI harms. And to the extent truly catastrophic scenarios keep us awake at night, those scenarios will be governed first by context-specific prospective legislation and regulation—just as aviation crashes and nuclear meltdowns are today—and only secondarily by ex-post compensation regimes.

The real question, then, is not whether tort can manage AI, but whether AI will one day eclipse tort entirely. If autonomous systems make serious accidents statistical rounding errors, the grand edifice of accident law—fault, causation, damages—will begin to empty. As Ryan Abbott put it, what role will remain for accident law once “no one dies from unintended injury”?³⁵⁵

That day is not yet here. Until it arrives, we ought to apply the tools we have, continue to refine them incrementally, and resist the siren song of special-liability regimes that promise certainty but deliver only cost. History urges—and rewards—patience. The electricity cases teach that tomorrow's mundane technology can feel, today, like a live wire. We should grasp it with the insulated implements already hanging on tort law's pegboard—and leave the breaker panel of strict liability alone.

³⁵⁵ Abbott, *supra* note 92, at 44.