
CONFRONTING AI LIABILITY AND INSURANCE[†]

KENNETH S. ABRAHAM^{*} & CATHERINE M. SHARKEY^{**}

CONTENTS

INTRODUCTION	1276
I. AI HARMS	1277
II. INSURANCE GAPS	1279
III. THE CALL FOR AFFIRMATIVE COVERAGE	1282
IV. AN EVOLUTIONARY VIEW OF AI LIABILITY AND INSURANCE.....	1285
CONCLUSION	1287

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

^{*} David and Mary Harrison Distinguished Professor of Law, University of Virginia School of Law.

^{**} Segal Family Professor of Regulatory Law and Policy, New York University School of Law.

INTRODUCTION

In *Artificial Intelligence, Judicial Evolution, and Insurance*, Professor Renee Henson achieves her stated goal of “advanc[ing] scholarly and industry dialogue on how best to confront the mounting challenges of algorithmic risk.”¹ We especially applaud her taking up our challenge to scholars to “recognize[] the interdependence of insurance and tort law development.”² Moreover, she has done so in an area worthy of sustained attention, namely emergent artificial intelligence (“AI”) harms.

The focus of Professor Henson’s article—“developing products liability lawsuits involving algorithmic tools”³—is, however, much narrower than the promise of its bold title. Professor Henson provides a very detailed analysis of a handful of state and federal court cases in order to illustrate a trend in doctrinal case law, whereby courts have shifted from denying tort claims of AI-related harms outright, primarily on Communications Decency Act Section 230 or First Amendment grounds, to recognizing emotional distress and mental health harms as “product”-related harms worthy of courts’ consideration.⁴

But, alas, paradoxically, this signature contribution of her article is at the same time its weakness and ultimately defeats any aspiration to provide a generalizable framework for consideration of the interplay between tort liability and insurance of AI harms. The crux of our critique is that Professor Henson’s subject matter designation of products liability cases and limitation of AI harms to emotional and mental harm skew both her diagnosis of the problem of insurance “gaps” and her proposed solution of accelerating the development of affirmative coverage.⁵ First, she writes little about the different causes of action (aside from products liability) that these kinds of

¹ Renee Henson, *Artificial Intelligence, Judicial Evolution, and Insurance*, 106 B.U. L. REV. 241, 304 (2026).

² *Id.* at 246, 246 n.12 (citing Kenneth S. Abraham & Catherine M. Sharkey, *The Glaring Gap in Tort Theory*, 133 YALE L.J. 2165, 2254-55 (2024) [hereinafter Abraham & Sharkey, *Glaring Gap*] (arguing actual and potential availability of insurance has influenced, and in our view should continue to influence, scope and nature of tort and other forms of civil liability)).

³ *Id.* at 253.

⁴ *See id.* at 245-46.

⁵ In a similar vein, we note the mismatch between Professor Henson’s rather broad claim that “traditional legal paradigms may be ill-equipped to resolve the novel legal questions presented by algorithmic liability,” *id.* at 265, and the narrow slice of evidence, in one distinct area of claims, upon which it is based. *Id.* at 246-47. Moreover, as we explore in a forthcoming article, while it is increasingly de rigueur for academics to claim that tort liability regimes are not up to the task of handling AI harms, such critiques are at best overstated and at worst misguided. *See* Kenneth S. Abraham & Catherine M. Sharkey, *Untangling AI Liability*, 115 CAL. L. REV. (forthcoming 2027) (manuscript at 2-3) (on file with authors) [hereinafter Abraham & Sharkey, *Untangling AI Liability*].

harms might generate, even assuming that the harms to which she limits herself are a coherent group for purposes of AI liability.⁶ Second, the category of “mental and emotional harms” is not wholly omitted from the coverage language of the insurance policies she addresses. Many policies provide silent coverage for AI harms, including some intangible losses.⁷ Third, significant moral hazard and adverse selection concerns jeopardize the feasibility of the new affirmative insurance she proposes.

We identify these limitations in the spirit of encouraging Professor Henson to carry this important quest even further. Moreover, we aim to convince her to embrace the alternative evolutionary view of the interaction between tort liability and insurance that we set out below.

I. AI HARMS

We recognize that sketching the full landscape of future liability for AI harms would be radically premature, given that we do not yet know at the necessary level of detail the range of different operations and applications that may involve AI. Professor Henson’s detailed survey of emerging products liability cases is a valuable contribution in its own right. Our point here, however, is that her exclusive focus on these cases and the mental and emotional harms associated with them necessarily limits her ability to generalize regarding the interplay between liability for AI and insurance.

Professor Henson describes in detail what she terms a “critical turning point” whereby “courts have shown a willingness to classify various platforms’ algorithmically driven tools as products for purposes of products liability” in cases that allege significant mental and emotional harms.⁸ While we agree this is a noteworthy trend—or even “a potentially seismic shift in

⁶ Notwithstanding our shared sense of the significance of products liability law to emergent AI harms, see Henson, *supra* note 1, at 248 nn.23 & 25 (citing Catherine M. Sharkey, *A Products Liability Framework for AI*, 25 COLUM. SCI. & TECH. L. REV. 240, 259 (2024)), we believe it is necessary to situate these developments within a wider framework of potential tort causes of action for a wide array of AI harms in order to make sense of the implications for AI liability and insurance. See Abraham & Sharkey, *Glaring Gap*, *supra* note 2, at 2169–71. Professor Henson instead seems to be assuming either that products liability law is set-in-stone or that some adjusted cause of action that resonates with products liability is infeasible. But why spend so much time on what the early courts addressing AI liability have said, and parsing what they have said so carefully, when this subject is so heavily *tabula rasa* and requires the same kinds of adjustments that the courts and Restatement made in the 1960s when the shift away from existing law was made?

⁷ See Kenneth S. Abraham & Catherine M. Sharkey, *Insurance and the Law of Artificial Intelligence*, 32 INS. L. REV. 1, 5–9 (2026) [hereinafter Abraham & Sharkey, *Insurance and the Law of AI*].

⁸ Henson, *supra* note 1, at 241 (arguing “[r]ecent jurisprudence . . . signals a critical turning point”); *id.* at 272.

judicial interpretations of social media platform liability”⁹—it still represents only a small slice of the universe of harms that AI may cause.

We would therefore situate Professor Henson’s analysis within a more general framework that acknowledges how conventional tort duties can be applied to AI harms for tortiously caused bodily injury, property damage, certain forms of emotional loss, economic loss caused by fraud, and the dignitary and emotional losses associated with “intentional” torts such as defamation and invasion of privacy.¹⁰ The first category involves certain forms of negligently caused physical harm associated with creating or using AI.¹¹ The second category consists of products liability for “cyber-physical” harms such as injury resulting from “self-driving” cars and medical AI.¹² The third and fourth categories involve liability for negligently inflicted emotional distress and pure economic losses, respectively.¹³ And the fifth consists of some of what are often called “intentional” torts, which includes defamation, invasions of privacy, tortious release of personal data, and fraud, among other torts.¹⁴ To be sure, there may be new forms of loss caused by AI that do not fit neatly into these conventional tort duties—or, as Professor Henson seems to imply, additional pressure on long-existing limitations on tort liability for freestanding emotional and mental harms¹⁵—but it is difficult to situate such a conclusion absent a more fulsome analysis of the range of AI harms across myriad tort causes of action.

Professor Henson draws upon a database of AI litigation, from which she selects the handful of products liability cases that she analyzes in detail in her article.¹⁶ Given that her selection criterion was cases “falling within the category of tort actions,”¹⁷ we especially do not understand (nor does she defend) her exclusion of “autonomous vehicle claims” or privacy cases,¹⁸ which would fit within our broader framework of causes of actions for AI harms (specifically, the second and fourth categories outlined above). Moreover, surely if Professor Henson is right that “recent decisions recognizing potential liability should set off seismic tremors that ought to

⁹ *Id.* at 276.

¹⁰ See Abraham & Sharkey, *Untangling AI Liability*, *supra* note 5, at 13-23.

¹¹ *Id.* at 14-15.

¹² *Id.* at 16-18.

¹³ *Id.* at 18-20.

¹⁴ *Id.* at 20-23. For elaboration, see *id.* Section III.B (arguing while AI’s opacity or “black box” nature will sometimes pose challenges for establishing tort liability, there are various of forms of tort liability which would not be impeded by AI’s opacity).

¹⁵ See Henson, *supra* note 1, at 285 (“As plaintiffs increasingly survive early motions to dismiss and allege emotional and mental harm, including self-injury, potential exposure expands.”).

¹⁶ *Id.* at 253 n.53.

¹⁷ *Id.*

¹⁸ *Id.*

alarm both defendants and insurers about the scale of future litigation and potential for exposure,”¹⁹ such concern would by no means be limited either to the realm of products liability suits or the domain of mental and emotional distress harms.

II. INSURANCE GAPS

We can usefully divide the AI liability insurance landscape into two categories: “silent” AI liability insurance and “affirmative” AI liability insurance. Silent AI liability insurance consists of coverage provided by existing types of liability insurance under which AI liability falls automatically into the general terms of coverage and from which AI liability is not excluded.²⁰ Affirmative AI liability insurance consists of express coverage of AI liability.²¹ AI is already covered by a number of existing forms of silent liability insurance, and we also predict the growth of affirmative AI coverage that expressly covers specified AI losses.²²

By contrast, Professor Henson takes the scope of silent coverage under existing forms of liability insurance as *de minimis* and static. From this baseline she then proposes new forms of affirmative insurance to address the gaps she identifies. According to Professor Henson, “[b]ecause algorithmic liability is expanding, insurance should—but currently does not—explicitly cover these risks.”²³ Specifically, she argues, “[t]his new liability exposes significant insurance gaps, as the associated harms—emotional, mental, and self-induced—are likely to fall outside the scope of coverage provided by existing commercial general liability (‘CGL’), businessowners’ policies (‘BOP’), errors and omissions (‘E&O’) or cyber insurance policies.”²⁴

Professor Henson’s identification of present insurance “gaps” follows directly from her truncated notion of AI harms as limited to “emotional and mental harm, including self-injury.”²⁵ Even in relation to these types of AI harms, however, she is overly dismissive of the existence of silent AI coverage. First, CGL insurance policies have a separate set of coverages, aside

¹⁹ *Id.* at 286.

²⁰ Abraham & Sharkey, *Insurance and the Law of AI*, *supra* note 7, at 4.

²¹ *Id.*

²² *See id.* at 5-9 (noting how widespread and varied forms of liability insurance already cover or would cover many different forms of AI liability).

²³ Henson, *supra* note 1, at 284.

²⁴ *Id.* at 245-46.

²⁵ *Id.* at 285; *see also id.* at 250-51, 254, 288-90, 303-05. Professor Henson makes an additional argument that courts would interpret such claims as falling outside of traditional CGL policies “because plaintiffs tend to allege that defendants intentionally and knowingly design defective products, even when claims are asserted under theories of negligence, the allegations sound in intentional misconduct.” *Id.* at 296. We are highly skeptical. Indeed, it is a rare design defect products liability claim that rises to the level of intentional misconduct.

from liability for bodily injury and property damage, termed “Personal and Advertising Injury Liability,”²⁶ that insures against liability for “defamation” and “invasion of privacy.”²⁷ This would cover liability for AI’s commission of or involvement in the torts of libel, slander, and false light, at a minimum.²⁸ Second, cyber insurance covers liability for certain data and network-related losses, as well as certain forms of liability for defamation and invasion of privacy.²⁹ It would cover these forms of liability when cyber-related harm arose out of AI. Both of these forms of liability insurance would cover AI liability under the circumstances that their terms of coverage already contemplate. To be sure, this is not all-encompassing AI liability insurance, but it shows that some of the things that could result in liability for emotional loss are covered by CGL and cyber policies.

²⁶ KENNETH S. ABRAHAM & DANIEL SCHWARZ, *INSURANCE LAW & REGULATION* 478 (8th ed. 2025).

²⁷ Kenneth S. Abraham, *Free Speech, Breathing Space, and Liability Insurance*, 111 VA. L. REV. 1007, 1016-17 (2025).

²⁸ Silent coverage of AI-induced defamation and invasion of privacy might also be provided by Media Liability insurance and/or Employment Practices Liability insurance, which covers liability for discrimination in hiring or in the terms of employment. *See id.* at 1023 (discussing Media Liability coverage for speech-related tort claims); *id.* at 1026 (explaining Employment Practices Liability insurance also often covers speech-related tort claims). Additional silent AI coverage may be provided by Tech Errors and Omissions and Malpractice Liability insurance, which cover losses resulting from the negligent provision of technical services and the provision of health care, respectively. *See, e.g.*, BEAZLEY, BEAZLEY MEDIA TECH, <https://www.beazley.com/globalassets/product-documents/policy-form/beazley-media-tech-policy-us.pdf> [<https://perma.cc/ZPS8-Q4D7>] (last visited Feb. 2, 2026) (covering liability for “failure of Tech Products to perform the function or serve the purpose intended”); *see also* Abraham & Sharkey *Insurance and the Law of AI*, *supra* note 7, at 6 (discussing medical malpractice insurance).

²⁹ A standard cyber liability insurance policy covers:

Electronic media peril means the display by you of electronic data on your internet or intranet site that directly results in any of the following:

1. Defamation, libel, slander, product disparagement or trade libel;
2. Invasion of an individual’s right of privacy or publicity, including false light, intrusion upon seclusion, commercial misappropriation of likeness, and public disclosure of private facts;
3. Plagiarism or misappropriation of ideas under an implied contract;
4. Infringement, misappropriation or violation of any copyright, trademark, trade name, trade dress, title, slogan, service mark or service name; or
5. Domain name infringement or improper deep-linking or framing.

PHILA. INDEM. INS. CO., CYBER SECURITY LIABILITY COVERAGE FORM, https://assets.ctfassets.net/ai8tc5m8qkn7/mYQ18qaukxtcGPrLCV5S/c6ef045b3ecb2c3d50d9125ee704ebdb/Cyber_Security_Liability_Policy_Form.pdf [<https://perma.cc/BG6C-2V9J>] (last visited Feb. 2, 2026).

Professor Henson mentions only in passing such silent AI coverage before positing that “insurers increasingly exclude silent coverage by explicitly stating that existing policies do not apply to AI-related injuries.”³⁰ Professor Henson summarily dismisses CGL policy coverage of Personal and Advertising Injury Liability, by reiterating that defamation and privacy claims fall outside “the algorithmic injury claims analyzed herein”—namely emotional distress and mental harms arising from products liability cases.³¹ In a similar fashion, she concludes “cyber insurance is also unlikely to address the coverage gaps” because “[a]lgorithmic injuries generally do not clearly implicate any of these covered categories.”³² But this leads us back (once again) to her artificial limitation of “algorithmic injuries”—here, in particular, her exclusion of privacy claims which would certainly implicate “unauthorized disclosure of personal information.”³³

The stakes for assessing the insurance “gaps” are twofold. First, Professor Henson claims that the gap in coverage “poses broader systemic concerns” including “threaten[ing] innovation in the AI sector, which is increasingly intertwined with matters of national security.”³⁴ More pointedly, Professor Henson claims that “[a] significant downturn in AI development due to liability exposure could jeopardize the United States’ competitive standing in the global AI landscape.”³⁵ Second, her omission of details regarding policy language and the scope of silent AI coverage relates directly to her call for “specialized insurance solutions,”³⁶ namely affirmative coverage, to which we now turn.

³⁰ Henson, *supra* note 1, at 298.

³¹ *Id.* at 292 n.306 (“Coverage B agreements are also unlikely to be triggered under standard policy language as none of the algorithmic injury claims analyzed herein clearly implicate ‘personal and advertising injur[ies],’ generally arising from slander or libel, violating one’s right to privacy, misusing another’s advertising idea, or infringing on another’s copyright.”).

³² *Id.* at 298.

³³ *Id.*

³⁴ *Id.* at 290.

³⁵ *Id.* at 291. Professor Henson asserts that “in the absence of current political and public will to shield firms from algorithmic liability, insurance may play a significant role in shaping the future trajectory of AI development.” *Id.* at 290. We are somewhat mystified by this assertion. Our own view is that tort liability for AI harms will have a modest impact upon AI innovation (much more modest than, say, stringent *ex ante* regulation); moreover, calls to immunize firms from liability altogether pose a much greater risk of under-deterrence. But, regardless, the existence of tort liability (and the strength of its deterrence signal) will undoubtedly shape the trajectory of AI development.

³⁶ *Id.* at 291.

III. THE CALL FOR AFFIRMATIVE COVERAGE

We are at the dawn of the development of affirmative coverage of AI liability. A few insurers offer niche coverage against liability risks that is described publicly in only rudimentary fashion.³⁷ There are already tiny bits of such insurance. Some of this new coverage appears to be for breach of performance guarantees.³⁸ Some covers general liabilities.³⁹ Some covers liability for data poisoning.⁴⁰ Some covers intellectual property liability.⁴¹ The scope and types of coverage will undoubtedly expand. A scenario we envision involves the spread of similarly targeted, discrete and limited areas of coverage, offered by increasing numbers of insurers, with both the scope and amount of coverage expanding over time, as claim and loss experience accumulates.⁴²

Professor Henson takes a more proactive position, advocating for the development of specialized algorithmic liability insurance to fill the existing “gaps”; “existing insurance frameworks do not cover the distinct risks posed by algorithmic liability, underscoring the need for affirmative coverage solutions.”⁴³ As mentioned above, her proposed solution is limited by her having framed the “distinct risks posed by algorithmic liability” overly narrowly, as pertaining solely to “mental and emotional distress and self-inflicted injury.”⁴⁴ She does not hold back her disdain of insurers who “with the issuance of limited affirmative AI policies . . . appear to cover AI-related harms in name only.”⁴⁵ But, here too, Professor Henson’s critique hinges on her cribbed definition of “AI-related harm”; indeed, Professor Henson has (to us, indefensibly) set aside from her analysis affirmative policies that “target economic harms—and potentially some physical harms—caused by technological malfunctions.”⁴⁶

But, even on her own terms, as she “seeks to hasten a market resolution,”⁴⁷ we expected to find a more robust discussion of the supply-side factors driving insurers’ incentives to offer coverage. First, as more insurers enter the market, competition may drive premium rates down and make it less attractive for more insurers to offer coverage. Second, insurers’ capacity to

³⁷ See *Insurance and the Law of AI*, *supra* note 7, at 10–11.

³⁸ *Id.* at 10.

³⁹ *Id.*

⁴⁰ *Id.* at 10–11.

⁴¹ *Id.* at 10.

⁴² See *id.* at 15–16.

⁴³ Henson, *supra* note 1, at 250; see also *id.* at 301.

⁴⁴ *Id.* at 250.

⁴⁵ *Id.* at 300–01.

⁴⁶ *Id.*

⁴⁷ *Id.* at 250; see also *id.* at 292 (“Existing insurance products are insufficient to bridge the coverage gaps inherent in algorithmic injury claims.”); *id.* Section IV.A (referencing section titled “Current Policies Have Coverage Gaps for Algorithmic Injury Claims”).

make underwriting decisions and engage in risk management of their potential and actual insureds may inhibit the development and expanded scope of coverage. The safety of different AI operations may initially be difficult for insurers to assess, and for a significant period of time insurers may be unable to identify or recommend effective loss prevention techniques. This is because AI safety will itself be comparatively opaque to insurers considering underwriting potential policyholders' applications.

As a consequence, the development of affirmative AI liability insurance may be hindered, periodically, by adverse selection and moral hazard challenges—the classic twin obstacles to the market-driven insurance function.⁴⁸ Adverse selection arises because insurers will be less able to assess the riskiness of prospective policyholders to the extent they can in other lines of insurance, as those who believe that they are at above-average vulnerability to suit are more likely to seek insurance against AI liability.⁴⁹ Further, once insured, policyholders may exercise less care than they would if they were not insured.⁵⁰ This moral hazard will be more difficult for insurers to neutralize than in other lines of insurance because insurers themselves will be incompletely informed about what steps policyholders can take to reduce their risk of loss. But the insurance industry has always overcome obstacles like this in its effort to capture business opportunities and thereby protect policyholders against new liability risks.⁵¹ Temporary troubles should not lead the courts and other policymakers to the conclusion that AI liability is not insurable.

Professor Henson does not wrestle adequately with these twin challenges facing insurers' provision of new affirmative AI liability coverage. First, Professor Henson discusses moral hazard solely in connection with plaintiffs' ability to feign emotional or mental injuries, and the difficulty of reliably verifying mental and emotional injuries, prompting limitations on recovery to avoid incentivizing frivolous litigation based on harms that are difficult to disprove.⁵² According to Henson, "[a]lthough these concerns are not without merit, this Article posits that they are overstated."⁵³ Setting to one side whether, as an empirical matter, this conclusion holds, the moral hazard associated with compensating for emotional harm afflicts the plaintiff (e.g., feigning emotional harm), not the defendant or the defendant's liability insurer. In any event, whatever moral hazard there would be from insuring against AI liability would also be present in any new, affirmative AI liability insurance that was developed. Given that it is not a problem that would afflict

⁴⁸ ABRAHAM & SCHWARCZ, *supra* note 26, at 6.

⁴⁹ *Id.* at 6-7.

⁵⁰ *Id.* at 8.

⁵¹ *See id.* at 8-10.

⁵² Henson, *supra* note 1, at 288-89; *see also id.* at 302-04.

⁵³ *Id.* at 289.

only existing insurance, it by no means can support an argument for a new form of insurance.

Moreover, adverse selection is a potential problem for new “affirmative” AI liability insurance, since only those who wanted that insurance would be applying for it, and some of those who wanted it would likely be at disproportionately above-average risk of suffering a loss but would be hard for insurers to identify as such.⁵⁴ In contrast, policyholders already buying CGL or cyber insurance would pose less risk of adverse selection because they would already have been buying this insurance or in any case would be buying much more than only AI insurance.

We find it impossible to evaluate Professor Henson’s proposed “market-based intervention” that “urges the insurance market to swiftly respond to this growing risk”⁵⁵ without undertaking some analysis of the moral hazard and adverse selection concerns as a potential disadvantage of the new insurance she proposes.

In our view, affirmative AI liability insurance will have two advantages worth considering when fashioning AI liability rules. First, the institution of insurance is already in place. As affirmative AI insurance spreads, it will do so smoothly. In contrast, other institutions—such as new regulatory agencies, if that approach is taken—will have to be built from the ground up. Second, although the devices that liability insurance employs to regulate the safety of its policyholders are very far from effectively achieving that goal, insurance can and does have some safety promoting effects.⁵⁶

We tend to agree with Professor Henson that “insurers are more capable and adaptable than often acknowledged.”⁵⁷ At the same time—and certainly without confronting the classic adverse selection and moral hazard concerns that plague any market-based insurance proposal—we do not share her blind faith that “[i]f insurers can devise mechanisms to cover risks such as maritime ship fraud, they should be able to develop models to insure against mental and emotional harms arising from algorithmic liability.”⁵⁸

⁵⁴ The phrase “adverse selection” does not appear in Professor Henson’s article, but she remarks that “[e]ssentially, insurers fear opening the flood gates by explicitly covering emotional and mental harm because these claims are difficult to define and limit.” *Id.* at 288 & n.291.

⁵⁵ *Id.* at 300.

⁵⁶ See Abraham & Sharkey, *Glaring Gap*, *supra* note 2, at 2237-41 (discussing variety of mechanisms by which liability insurance has reduced incidence of tortiously caused losses).

⁵⁷ Henson, *supra* note 1, at 303.

⁵⁸ *Id.* at 304.

IV. AN EVOLUTIONARY VIEW OF AI LIABILITY AND INSURANCE

We take an evolutionary view of the interaction between tort liability and insurance. Liability insurance policies are frequently revised so as to provide more coverage, when there is demand for it, or to exclude coverage, when a new form of liability arises.⁵⁹ Then, as experience accumulates, insurers either relax the exclusions so as to re-include the coverage in question, or develop new forms of targeted, specialized insurance to provide the now-excluded coverage.⁶⁰

We therefore expect that, as CGL and other liability insurers begin to receive claims, they will develop exclusions that preclude certain forms of AI-related liability that they conclude their policies are not suited to cover. A few insurers already are including AI exclusions in certain policies.⁶¹ If and when that continues, it will help create what may by then be increasing demand for coverage against AI liability and thereby catalyze the growth of affirmative coverage.

We can look to the evolutionary development of cyber insurance as a guide. First, cyber-related liability covering physical harm was covered because it was not expressly excluded; then it was excluded and a wholly new form of specialized cyber insurance was developed.⁶² Professor Henson likewise sees trends on AI coverage that “mirror[] the early development of cyber insurance.”⁶³ But whereas Professor Henson sees “high premiums, narrow coverage, and numerous exclusions” leading her to “characterize it as largely illusory,”⁶⁴ we see that today, cyber insurance is a thriving, \$16.6 billion business in the United States alone, up from just \$2 billion in 2018.⁶⁵

Professor Henson may well be correct that “[a]s algorithmic liabilities continue to expand, it is plausible that insurers will issue exclusions under CGL and related policies specifically targeting . . . algorithmic injuries.”⁶⁶ But—as the unfolding story of cyber insurance demonstrates—even if that is what is happening right now, that is not necessarily the end of the story. The

⁵⁹ See Abraham & Sharkey, *Glaring Gap*, *supra* note 2, at 2251-52 (discussing evolution of insurance for cyber liability).

⁶⁰ See *id.*

⁶¹ See Geoffrey B. Fehling, Michael S. Levine, Madalyn Moore, Alex D. Pappas & Hunton Andrews Kurth, *The Continued Proliferation of AI Exclusions*, NAT'L L. REV. (May 28, 2025), <https://natlawreview.com/article/continued-proliferation-ai-exclusions> [<https://perma.cc/HN-M5-EVYK>] (discussing insurance company's exclusion of AI harms from liability coverage).

⁶² See Abraham & Sharkey, *Glaring Gap*, *supra* note 2, at 2251-52.

⁶³ Henson, *supra* note 1, at 300.

⁶⁴ *Id.*

⁶⁵ For elaboration, see Abraham & Sharkey, *Insurance and the Law of AI*, *supra* note 7, at 4.

⁶⁶ Henson, *supra* note 1, at 299.

exclusions might serve as a catalyst for the development of targeted specialized affirmative coverage.

Even that scenario is not assured. It is worth noting that one of the challenges associated with the development of affirmative liability insurance is defining what is and is not covered. In this case, insurers will have to define “AI” and “AI Liability,” or some subset of it. In contrast, the existing forms of insurance, which as to AI are silent rather than affirmative, would have no need to do that if insurers were to decide eventually that they wanted to cover it. It may well be that as AI liability develops, there are a lot of “border” disputes between affirmative AI insurers and their policyholders about what falls within and what falls outside the policies’ definitions of what they cover. Consequently, it may prove more efficient for existing policies simply not to exclude coverage of AI liability, to price their silent coverage into their premiums, and to cover it. If that occurred, then affirmative AI insurance would either not be needed and would fade away or would evolve into a niche product that did not do what Professor Henson envisions it would.

We simply do not know which way things will go in the long run. But, at this time, we can at least recognize the dilemma and consider possible alternative scenarios. A fulsome analysis should consider the prospect that first-party insurance for the victims of AI-related losses will become available. Existing forms of first-party insurance, including property, cyber, and business interruption insurance, may provide silent coverage of some of the losses that result from AI. New forms of affirmative first-party insurance against AI losses may also develop. Minimal amounts and types of such insurance already exist;⁶⁷ over time, they are likely to grow.

On the one hand, such insurance could weaken the case for imposing tort liability for these losses. On the other hand, first-party insurance could actually provide victims with the financial “staying power” to support their lawsuits against those who have caused these losses.⁶⁸ Thus, sometimes the existence and availability of first-party insurance for those who experience particular forms of loss serves as an argument against imposing liability on parties who cause such losses.⁶⁹ But in the AI context, the silent coverage

⁶⁷ See Abraham & Sharkey, *Insurance and the Law of AI*, *supra* note 7, at 10-11.

⁶⁸ Professor Henson relegates discussion of first-party insurance to a footnote where she argues it is “unlikely to attract a sufficiently broad customer base . . . [and] premium costs may be prohibitively high, particularly in the initial stages when such coverage would constitute a luxury-line offering.” Henson, *supra* note 1, at 300 n.354. Once again, Professor Henson’s lens is extremely constrained, focused entirely on the limited coverage for mental and emotional harms that remains her exclusive focus.

⁶⁹ The “stranger” economic loss rule in torts, for example, often precludes imposing tort liability for economic loss suffered by third parties whose losses result from a tortfeasor’s negligently causing bodily injury or property damage to another party. See Catherine M. Sharkey, *Can Data Breach Claims Survive the Economic Loss Rule*, 66 DEPAUL L. REV. 339, 341 (2017). Thus, a driver whose negligence causes a collision on a bridge is not liable for the economic losses suffered by businesses on the other side of the bridge that result from

provided by any existing first-party insurance, or any affirmative first-party insurance against AI-related losses that develops, might instead be considered consistent with the development of rules favoring tort liability for these losses.

CONCLUSION

We join ranks with Professor Henson in recognizing that “the onset of credible legal risk prompts questions regarding when and if insurance markets will ultimately respond”⁷⁰ and, moreover, share her view that “the insurance market is capable of insuring algorithmic liability.”⁷¹

We part ways by insisting that AI harms are wide-ranging and by no means limited to the sphere of “emotional, mental, and self-induced.”⁷² We also take a more evolutionary view of how AI liability and the insurance market will interact. At present, there is a considerable amount of existing insurance that will cover AI liability, and the future prospects for development of targeted AI liability insurance are good. We would thus urge a modicum of patience in the face of so-called coverage “gaps.” Moreover, given the twin forces of adverse selection and moral hazard, we are not ready to sign onto Professor Henson’s new affirmative coverage proposal.

potential customers’ lack of access to these businesses while the bridge is closed. These businesses’ access to Contingent Business Interruption insurance is sometimes cited as part of the rationale for the economic loss rule. *See, e.g., id.* at 352 n.51.

⁷⁰ Henson, *supra* note 1, at 287.

⁷¹ *Id.* at 249.

⁷² *Id.* at 245.