

TEMPLE LAW REVIEW

© 2024 TEMPLE UNIVERSITY OF THE COMMONWEALTH SYSTEM OF
HIGHER EDUCATION

VOL. 96 NO. 3

SPRING 2024

ARTICLES

“I AM BECOME DEATH, THE DESTROYER OF WORLDS”: APPLYING STRICT LIABILITY TO ARTIFICIAL INTELLIGENCE AS AN ABNORMALLY DANGEROUS ACTIVITY

*Renee Henson**

Artificial intelligence (AI)-enabled tools have produced a myriad of injuries, up to and including death. This burgeoning technology has caused scholars to ask questions, such as, How do we create a legal framework for AI? Because AI creators have acknowledged that even they do not know the capacities of their technology for good or bad outcomes, this Article argues that an existing framework, strict liability, is an appropriate fit for harms arising from this new technology because a party need not prove negligence to prevail. Strict liability was uniquely developed to handle those activities that are “abnormally dangerous.” An abnormally dangerous activity is one that imposes an abnormal risk on anyone who is in the vicinity of its use. The quintessential historical example of this is strict liability applied to the production of atomic energy. Congress acknowledged that nuclear energy would be extremely beneficial to society but could not be supported by the safety net of insurance, due to the potentially catastrophic results from its production. Congress enacted the Price-Anderson Act to both establish insurance for nuclear plant operators and to set a

* Visiting Assistant Professor of Law at the University of Missouri School of Law. The author previously practiced as a business and commercial litigation attorney at Stinson LLP, representing clients engaged in complex matters including trademark and copyright infringement and products liability lawsuits. She earned a J.D. at University of Missouri School of Law and served as the Editor in Chief of the Business, Entrepreneurship & Tax Law Review. She earned a B.A. at Columbia College magna cum laude with Distinction.

The author expresses her utmost gratitude, love, and respect to her husband, Professor Chuck Henson, for his everlasting support and encouragement. She expresses sincere thanks to Nicole McConlogue, Sandra Sperino, Vanessa Browne Barbour, Gary Meyers, Dennis Crouch, and Catherine Smith for their comments in support of this Article.

liability cap. The Act served as a carrot to encourage nuclear operator entrepreneurs and as a protection for the public. The development of nuclear energy is comparable to the development of AI. Nuclear energy and AI share the essential feature that their creators acknowledge the potentially enormous, but not fully understood, capacities of their creations to do harm. This Article begins by discussing the development of strict liability for emerging technologies with the attribute of being “abnormally dangerous.” It then explores the issues associated with applying a strict liability framework to AI and posits that an umbrella insurance protection similar to the Price-Anderson Act would be a viable solution to one of the most salient questions in modern history: How do we create a legal framework for AI? This Article argues that regulation should create a compensatory structure for potentially catastrophic harms created by an unknown (or not fully understood) technology.

INTRODUCTION..... 351

I. AN OVERVIEW OF ARTIFICIAL INTELLIGENCE, PUBLIC PERCEPTION, AND GOVERNMENT CONCERNS..... 355

 A. *What Is AI?*..... 355

 B. *Types of Predictive Algorithms* 356

 C. *Rule-Based Algorithms*..... 356

 D. *Machine Learning AI*..... 357

 E. *There Is Widespread Concern that Artificial Intelligence Will Introduce Risks of Injury that Are Unlike Anything Seen in History*..... 358

II. STRICT LIABILITY IS APPROPRIATE FOR THE DANGERS ASSOCIATED WITH AI BECAUSE THE USE OF AI MAY BE CONSIDERED AN “ABNORMALLY DANGEROUS ACTIVITY” 362

 A. *Certain AI Uses May Satisfy the Abnormally Dangerous Activities Test*..... 363

 1. *The High-Risk Management Tool is an AI-Enabled Activity* 364

 B. *Applying the Six-Factor Test to the Health-Management Tool*.... 366

 1. *Inability To Eliminate the Risk by Exercising Reasonable Care*..... 366

 2. *Existence of a High Degree of Risk of Some Harm to the Person, Land, or Chattels of Another*..... 370

 3. *Likelihood that the Harm that Results Will Be Great* 371

 4. *Extent to Which the Activity Is Not a Matter of Common Usage* 372

 5. *Inappropriateness of the Activity to the Place Where It Is Carried On*..... 374

 6. *The Extent to Which Its Value to the Community Is Outweighed by Its Dangerous Attributes* 375

C.	<i>On Balance, the High-Risk Management Tool May Be Considered Abnormally Dangerous When Analyzed Under the Abnormally Dangerous Activities Test</i>	376
III.	THE SIX-FACTOR ABNORMALLY DANGEROUS ACTIVITIES TEST IS USEFUL, BUT SHOULD BE REVISED TO FIT THE SPECIAL CHALLENGES ASSOCIATED WITH ARTIFICIAL INTELLIGENCE.....	376
A.	<i>The Inability To Foresee Harms Associated with Its Use, Even When Reasonable Care Is Exercised</i>	376
1.	<i>At the Time of Its Emergence, Nuclear Energy Was Also Deemed Unpredictable, and Therefore Abnormally Dangerous</i>	380
B.	<i>Eliminating the “Extent to Which the Activity Is Not a Matter of Common Usage” Requirement</i>	381
C.	<i>Eliminating the “Inappropriateness of the Activity to the Place Where It Is Carried On” Requirement</i>	382
IV.	UMBRELLA PROTECTIONS CONSISTENT WITH THE PRICE-ANDERSON ACT WOULD SERVE TO SHIELD CONSUMERS WHILE ENCOURAGING THE AI INDUSTRY TO FLOURISH.....	383
A.	<i>The Concept of Nuclear Energy Insurance Policies Can Be Applied to the Harms Associated With Artificial Intelligence</i>	384
1.	<i>A Two-Tiered Insurance Approach Protects the Consuming Public While Limiting Liability</i>	387
B.	<i>A Combined Strict Liability and Two-Tiered Insurance Policy Approach Would Address Concerns Regarding AI Regulation</i>	389
CONCLUSION		390

INTRODUCTION

Artificial intelligence (AI) can discover antibiotics, plan wars, and diagnose diseases. AI-enabled tools are here to stay and are only increasing in number. That said, people are fearful of AI. “How Could AI Destroy Humanity?”—it is almost daily that the news media issues breathless articles like this to warn the public about the harms that AI will do.¹ Stephen Hawking, Sam Harris, and Elon Musk have all warned that AI may present an existential threat to humanity.² The risks lead to an important question: *How do we create a legal framework for AI?* That is, how do regulators—courts and legislatures—create a framework to compensate victims injured by AI while not limiting

1. See, e.g., Cade Metz, *How Could A.I. Destroy Humanity?*, N.Y. TIMES (June 10, 2023), <https://www.nytimes.com/2023/06/10/technology/ai-humanity.html> [https://perma.cc/7Y38-3MYS].

2. MIRJANA STANKOVIC, RAVI GUPTA, BERTRAND ANDRE ROSSERT, GORDON I. MYERS & MARCO NICOLI, EXPLORING LEGAL, ETHICAL AND POLICY IMPLICATIONS OF ARTIFICIAL INTELLIGENCE 5 (2017), https://www.researchgate.net/publication/320826467_Exploring_Legal_Ethical_and_Policy_Implications_of_Artificial_Intelligence [https://perma.cc/E8CN-ZX32]; Clare Duffy & Ramishah Maruf, *Elon Musk Warns AI Could Cause ‘Civilization Destruction’ Even As He Invests in It*, CNN (Apr. 17, 2023, 9:35 PM), <https://www.cnn.com/2023/04/17/tech/elon-musk-ai-warning-tucker-carlson/index.html> [https://perma.cc/2M5C-2KDH].

technological entrepreneurship to ensure that the United States continues to be recognized as a cutting-edge innovator in the global market?³

The father of the nuclear bomb, Robert Oppenheimer, recalled a quote from the *Bhagavad-Gita* when he famously said, “Now, I am become Death, the destroyer of worlds,” in recounting his experience watching the explosion of the first nuclear bomb.⁴ The nuclear bomb was a dangerous new technology that Oppenheimer lost control of just as soon as it was developed. Following the first use of the nuclear bomb, in Hiroshima, Oppenheimer questioned the morality of the technology and expressed his concerns to President Truman and Congress.⁵ Oppenheimer was conflicted by his desire to “not hold back progress” and to ensure that the atomic bomb “would be a ‘hope’ and not a ‘peril’” to society.⁶ Oppenheimer felt responsible for the risks that the atomic bomb created.⁷ Oppenheimer presciently stated, “[I]t is my judgment in these things that when you see something sweet, you go ahead and do it and you argue about what to do about it only after you have had your technical success.”⁸ Indeed, the job of “what to do about it” fell to the United States government.

Congress later addressed some of the risks associated with nuclear energy by creating the Atomic Energy Act in 1954, which permitted nongovernment entities to use nuclear energy for the first time.⁹ Congress amended the Atomic Energy Act in 1957, creating the Price-Anderson Act.¹⁰ The Price-Anderson Act requires nuclear operators to be licensed by the Nuclear Regulatory Commission and requires each operator to purchase insurance to protect against the harms that would result from a nuclear

3. See Clark D. Asay, *Artificial Stupidity*, 61 WM. & MARY L. REV. 1187, 1193 (2020); *Software and AI as a Medical Device Change Programme Roadmap*, GOV.UK (June 14, 2023), <https://www.gov.uk/government/publications/software-and-ai-as-a-medical-device-change-programme/software-and-ai-as-a-medical-device-change-programme-roadmap> [<https://perma.cc/6AAF-RZZX>].

4. NBCUniversal Archives, *Atomic Bombings of Hiroshima and Nagasaki*, YOUTUBE, at 1:21 (Aug. 3, 2015), <https://www.youtube.com/watch?v=cY8q1ky3dLY> [<https://perma.cc/3CMS-DWH7>] (“We knew the world would not be the same . . . I remembered the line from the Hindu scripture, the Bhagavad Gita; Vishnu is trying to persuade the prince that he should do his duty, and to impress him, takes on his multiarmed form and says, ‘Now I am become Death, the destroyer of worlds.’ I suppose we all thought that, one way or another.”).

5. KAI BIRD & MARTIN J. SHERWIN, *AMERICAN PROMETHEUS: THE TRIUMPH AND TRAGEDY OF J. ROBERT OPPENHEIMER* 331–32 (Vintage Books 2006) (2005).

6. Barton J. Bernstein, *The Oppenheimer Loyalty-Security Case Reconsidered*, 42 STAN. L. REV. 1383, 1395–96 (1990) (first quoting CARLSBAD CURRENT-ARGUS, Aug. 17, 1945, at 1, col. 4; and then quoting ROBERT OPPENHEIMER, *Address to the Association of Los Alamos Scientists*, in ROBERT OPPENHEIMER: LETTERS AND RECOLLECTIONS 316–18 (Alice Kimball Smith & Charles Weiner eds., 1980)).

7. Oppenheimer famously told President Truman in a visit to the White House, “I feel I have blood on my hands.” Truman responded to Oppenheimer saying, “[T]he blood [is] on my hands . . . [so] let me worry about that.” BIRD & SHERWIN, *supra* note 5, at 332.

8. Lawrence C. Marshall, *Intellectual Feasts and Intellectual Responsibility*, 84 NW. U. L. REV. 832, 840 n.42 (1990) (quoting 2 U.S. ATOMIC ENERGY COMM’N, IN THE MATTER OF J. ROBERT OPPENHEIMER 266 (Apr. 12–13, 1954)). Professor Marshall slightly misquotes Oppenheimer, as the transcript reads “when you see something *technically* sweet.” U.S. ATOMIC ENERGY COMM’N, *supra* (emphasis added).

9. Atomic Energy Act of 1954, Pub. L. No. 83-703, 68 Stat. 919 (codified as amended at scattered sections of 42 U.S.C.); William D. O’Connell, Note, *Causation’s Nuclear Future: Applying Proportional Liability to the Price-Anderson Act*, 64 DUKE L.J. 333, 335 (2014).

10. Atomic Energy Damages Act (Price-Anderson Act), Pub. L. No. 85-256, 71 Stat. 576 (1957) (current version at 42 U.S.C. § 2210); see also O’Connell, *supra* note 9, at 335.

incident.¹¹ The Price-Anderson Act also requires nuclear operators to have a second layer of insurance.¹² Such a requirement ensures that if a nuclear incident creates more harm than the first layer can cover, each licensed nuclear operator must pay a pro rata charge, up to a certain amount to cover the costs of the harm.¹³

Built into the Price-Anderson Act is the requirement that all nuclear operators waive their defenses for any event determined to be an “extraordinary nuclear occurrence” under the Act.¹⁴ Plaintiffs must only prove causation and damages for a qualifying nuclear event.¹⁵ The Price-Anderson Act’s “no-fault” structure is similar to the strict liability structure, where a plaintiff must only prove causation and damages for qualifying causes of action.¹⁶ Oppenheimer and the Price-Anderson Act provide a warning and a solution, respectively; both are applicable to AI.

This Article argues that courts and legislatures do not need to *wholly* reinvent the legal wheel to address the problem of a beneficial technology with unknown but potentially limitless capacity to do harm. Strict liability is a type of “no-fault” liability that was developed, in part, for the most abnormally dangerous and ultrahazardous technological advance in history—atomic energy.¹⁷ Section I of this Article generally describes AI and the widespread concern that American society has regarding its potential harms. In Section II, the Article argues that strict liability is the most appropriate legal framework for the harms associated with AI, because a plaintiff does not need to prove negligence to prevail.¹⁸ This removal of the requirement to show negligence is important in the AI context for two reasons: (1) it would be difficult to establish a defendant’s duty based on its conduct because the general zone of foreseeable danger regarding AI use is murky, and (2) AI’s unpredictability makes foreseeability almost impossible to predict.

Historically, strict liability has been applied to activities that are abnormally dangerous, *inter alia*.¹⁹ This Article asserts that there are some types of AI-enabled

11. 42 U.S.C. § 2210(a); O’Connell, *supra* note 9, at 335.

12. § 2210(b)(1).

13. *Id.* § 2210(b)(1)(E); O’Connell, *supra* note 9, at 335; *see also infra* Part IV.A.

14. § 2210(n)(1); O’Connell, *supra* note 9, at 363 n.214; *see also infra* note 260 for the definition of an “extraordinary nuclear occurrence.”

15. *See* S. REP. NO. 89-1605, at 9 (1966) (“Suffice it to say at this point that, generally speaking, it is intended that the effect of these waivers will be to require a victim of an extraordinary nuclear occurrence, as that term is defined in the bill, to prove only that he or his property was damaged and that such damage was caused by the nuclear incident. Such waivers would be incorporated in [Atomic Energy Commission’s] indemnity agreements and in insurance policies and contracts which are required by the AEC to be furnished as proof of financial protection, and under mandate of Federal statute would be judicially enforceable in accordance with their terms.”).

16. *Id.*

17. *See In re Hanford Nuclear Rsr. Litig.*, 534 F.3d 986, 1004–06 (9th Cir. 2008); *see also* *Seay v. Chrysler Corp.*, 609 P.2d 1382, 1384 (Wash. 1980); Price-Anderson Amendments Act of 1988, Pub. L. No. 100-408, 102 Stat. 1066 (codified as amended in scattered sections of 42 U.S.C.), *reprinted in* 1 LEGISLATIVE HISTORY OF THE PRICE-ANDERSON AMENDMENTS ACT OF 1988, at 8 (1988).

18. *See* Ralph C. Anzivino, *The Implied Warranty of Merchantability and the Remote Manufacturer*, 101 MARQ. L. REV. 505, 512 (2017).

19. Stephen A. Evans, Comment, *Using the Abnormally Dangerous Activity Doctrine To Hold Principals Vicariously Liable for the Acts of Toll Manufacturers*, 21 B.C. ENV’T AFFS. L. REV. 587, 602–03 (1994) (“The

activities that would pass the “abnormally dangerous activities” test that courts have developed to determine whether strict liability is applicable to the harm:

[1] the existence of a high degree of risk of some harm to the person, land or chattels of others; [2] likelihood that the harm that results from it will be great; [3] inability to eliminate the risk by the exercise of reasonable care; [4] extent to which the activity is not a matter of common usage; [5] inappropriateness of the activity to the place where it is carried on; and [6] extent to which its value to the community is outweighed by its dangerous attributes.²⁰

This Article applies these six factors to a widely used AI predictive tool that uses patient data to determine high-risk care management.²¹ The use of this tool leads to disastrous results concerning its preference to choose White patients over Black patients as being more deserving of extra care. Arguably, this AI tool passes the abnormally dangerous activities six-factor test.²²

Although strict liability is appropriate for some uses of AI, the six-factor test used to distinguish between negligence and strict liability is not fully satisfactory in the AI context. Simply put, problems with AI are inherent in this extremely unique technology. Thus, although the abnormally dangerous activities test is useful, Section III of this Article proposes changes to the six-factor test to make it a more practical tool that has the flexibility to accommodate AI’s major acknowledged deficiency: its creators do not know precisely what it is doing, or how it arrives at its conclusions.

Section IV of this Article argues that strict liability is an appropriate legal framework for some harms associated with the use of AI and that an umbrella-insurance paradigm similar to that of the Price-Anderson Act would practically answer the question: *How do we create a legal framework for AI?* There are many potential solutions to this issue. This Article seeks to propose a tort model and regulatory structure to quickly compensate victims. The amended strict liability approach, combined with the mandatory two-tiered insurance requirement, would lead to greater consumer protections and would encourage the growth of the AI industry by limiting its liability while at the same time creating accountability for AI companies.

This Article posits that the law should not continue to permit the AI industry to externalize the costs of injuries. It is widely acknowledged that AI’s benefits to consumers are enumerable and have exceeded technology futurists’ original expectations.²³ At the same time, allowing AI businesses to continue to profit without legal accountability, or with only toothless legal accountability for the harms to persons or property, is undesirable and threatens consumer safety. To address problems that have

modern doctrine of strict liability for abnormally dangerous activities developed from the English case of *Rylands v. Fletcher*, decided in 1868.” (footnote omitted)).

20. RESTATEMENT (SECOND) OF TORTS § 520 (AM. L. INST. 1977).

21. See *infra* Part II.A.1.

22. This Article does not argue strict liability should be applied to *all* harms associated with injuries resultant from AI use. Instead, strict liability is a useful starting point that must be further developed and is only applicable when the injury that results is to persons, land, or personal property, and where the AI-enabled tool is one that might be considered an abnormally dangerous activity in its use.

23. See Emerging Tech. from the arXiv, *Experts Predict When Artificial Intelligence Will Exceed Human Performance*, MIT TECH. REV. (May 31, 2017), <https://www.technologyreview.com/2017/05/31/151461/experts-predict-when-artificial-intelligence-will-exceed-human-performance/> [https://perma.cc/2UQL-PX4H].

not existed until now and to answer the question “How do we create a legal framework for AI?” we must borrow from well-established doctrine while also innovating.

I. AN OVERVIEW OF ARTIFICIAL INTELLIGENCE, PUBLIC PERCEPTION, AND GOVERNMENT CONCERNS

A. *What Is AI?*

AI is not easily defined. AI may be summarized as a system that predicts, recommends, or decides outcomes that influence environments.²⁴ AI uses human and machine data inputs to “perceive real and virtual environments; . . . abstract such perceptions into models through analysis in an automated manner; and . . . use model inference[s] to formulate options for information or action.”²⁵ In simpler terms, AI consists of systems that permit machines to do things that would ordinarily require human intelligence to complete.²⁶ There are various types of AI that use different methodologies. Though a complete exploration of AI is beyond the scope of this Article, it will discuss people’s perceptions of AI, several types of AI algorithms, and the types of AI most relevant to this Article—rule-based and machine learning AI.

AI advances are moving faster than the general public can keep up with.²⁷ Mostly, the public’s response to AI has been to shun it, mock it, or display resignation about its capabilities.²⁸ The ambivalence about AI may be a fear-based response.²⁹ For example, some lawyers are concerned with AI’s ability to master “language fluency,” which could result in attorney job loss.³⁰ The fear of AI taking over legal services may unconsciously lead lawyers to underestimate AI’s true capabilities. On the other end of the spectrum, some people conceptualize AI in a magical way, believing that it is infallible.³¹ The bottom line is that the public is afraid of and confused by AI.³² Its growth and threatening lore seem to be leading to greater confusion.³³

24. 1 RAYMOND T. NIMMER, JEFF C. DODD & LORIN BRENNAN, INFORMATION LAW § 1:14 (2023).

25. *Id.*

26. *See id.*

27. *See* Tatum Hunter, *3 Things Everyone’s Getting Wrong About AI*, WASH. POST (Mar. 30, 2023, 7:27 AM), <https://www.washingtonpost.com/technology/2023/03/22/ai-red-flags-misinformation/> [<https://perma.cc/4G2N-8QDX>].

28. *See* Max Tegmark, *The ‘Don’t Look Up’ Thinking that Could Doom Us with AI*, TIME (Apr. 25, 2023, 6:00 AM), <https://time.com/6273743/thinking-that-could-doom-us-with-ai/> [<https://perma.cc/7JFA-NY9D>].

29. *See* Shep Hyken, *Half of People Who Encounter Artificial Intelligence Don’t Even Realize It*, FORBES (June 10, 2017, 9:12 AM), <https://www.forbes.com/sites/shephyken/2017/06/10/half-of-people-who-encounter-artificial-intelligence-dont-even-realize-it/?sh=788ecaf2745f> [<https://perma.cc/EL6C-9Q3F>].

30. Steve Lohr, *A.I. Is Coming for Lawyers, Again*, N.Y. TIMES (Apr. 10, 2023), <https://www.nytimes.com/2023/04/10/technology/ai-is-coming-for-lawyers-again.html> [<https://perma.cc/G3GK-LASX>].

31. *See* Hunter, *supra* note 27.

32. *See* Paul Ford, *Our Fear of Artificial Intelligence*, MIT TECH. REV. (Feb. 11, 2015), <https://www.technologyreview.com/2015/02/11/169210/our-fear-of-artificial-intelligence/> [<https://perma.cc/J9X6-BAHW>]; Hunter, *supra* note 27.

33. *See* Hunter, *supra* note 27.

B. *Types of Predictive Algorithms*

An algorithm, at its most basic level, is a computational process for solving a problem or accomplishing an end goal.³⁴ A prominent example application is Google's proprietary algorithm that ranks websites based on keyword searches.³⁵ The large-scale implementation of predictive algorithms has been aided by the growth of computing power, which permits analysis of a tremendous amount of data.³⁶

There are two ways that algorithms may be used in a predictive manner: (1) actuarially and (2) clinically.³⁷ Actuarial-based predictive decisions use correlations between variables and outcomes.³⁸ An example of actuarial-based predictive models is an AI system that diagnosis cancer. The result is not based on a specific person or set of circumstances, it is instead based on large data sets, which point to a particular outcome.³⁹

On the other hand, clinical predictions are used to outsource many complex societal decisions.⁴⁰ For example, AI systems that conduct clinical predictions are used to determine the likelihood that a child will be abused if left with their parents, the likelihood that a college student will accept an offer to a college based on a specific scholarship amount, and the likelihood that a prisoner will commit another offense if paroled.⁴¹ Given the uniqueness of these situations and the multitude of variables, the predictions are open-ended in nature—they do not provide a specific directive, but give the decision-maker additional information to consider.⁴²

C. *Rule-Based Algorithms*

The most familiar AI tools are rule-based models that determine outcomes using a series of preset rules.⁴³ Rule-based AI is structured in an "if x, then y" schema.⁴⁴ If certain prescribed conditions are present, then the AI will take the prescribed action or reach a

34. *Algorithm*, MERRIAM-WEBSTER'S DICTIONARY, <https://www.merriam-webster.com/dictionary/algorithm> [<https://perma.cc/FHV8-D97A>] (last visited Apr. 5, 2024) (defining algorithm broadly as "a step-by-step procedure for solving a problem or accomplishing some end").

35. See Julian Wallis, *How Does Google Search Work? Google's Search Algorithm Explained*, INTUJI: TECH BEHIND (Feb. 3, 2023), <https://intuji.com/how-does-google-search-work/> [<https://perma.cc/LYU6-8E5W>].

36. Robert Brauneis & Ellen P. Goodman, *Algorithmic Transparency for the Smart City*, 20 YALE J.L. & TECH. 103, 111 (2018).

37. See *id.*

38. *Id.* at 112.

39. See *id.*; see also FREDERICK SCHAUER, PROFILES, PROBABILITIES, AND STEREOTYPES 19–22 (2003).

40. See Brauneis & Goodman, *supra* note 36, at 111.

41. *Id.* at 111–12.

42. See *id.* at 112; see also William M. Grove, David H. Zald, Boyd S. Lebow, Beth E. Snitz & Chad Nelson, *Clinical Versus Mechanical Prediction: A Meta-Analysis*, 12 PSYCH. ASSESSMENT 19, 19 (2000).

43. See Yavar Bathace, *The Artificial Intelligence Black Box and the Failure of Intent and Causation*, 31 HARV. J.L. & TECH. 889, 898 (2018); Gina-Gail S. Fletcher, *Deterring Algorithmic Manipulation*, 74 VAND. L. REV. 259, 287 (2021).

44. See Edwina L. Rissland, *Artificial Intelligence and Law: Stepping Stones to a Model of Legal Reasoning*, 99 YALE L.J. 1957, 1965 (1990).

particular conclusion.⁴⁵ The rules are determined by AI designers and are applied to a set of facts.⁴⁶ The rule-based model is deterministic, rather than probabilistic.⁴⁷ Rule-based models do not change after their initial design.⁴⁸

D. Machine Learning AI

There is another category of AI that determines its outcomes based on what it learns from data sets that are supplied to it—machine learning.⁴⁹ In general, machine learning AI is trained on historical data sets and is then tested on new data sets to determine the machine learning’s validity, depending on what its designers instructed it to do.⁵⁰ Machine learning may change after its initial design.⁵¹ Machine learning models are complex and, in some cases, impossible to understand, because they are built with deep neural networks that act as webs of interconnected layers of electrical pathways that identify patterns in data that can be indeterminable.⁵² A neural network, as its name suggests, mimics the human brain where neurons send electrical signals via connecting synapses, which then trigger other neurons in a chain reaction.⁵³ Neural networks arise when there are multiple layers of neurons.⁵⁴ As it relates to AI, each “neuron” consists of a mathematical function that performs an individual specified task.⁵⁵ These neural networks work as AI’s “internal engine.”⁵⁶

Another category of machine learning AI is unknowable to humans because it “thinks” by locating “geometric patterns among those variables that humans cannot visualize”—meaning, in effect, that it “sees” dimensions that human beings cannot perceive.⁵⁷ This unknowable quality illustrates what is commonly referred to as the “black box problem”—that even AI’s human creators cannot understand or predict how it comes to the decisions it produces.⁵⁸ As discussed in greater detail in Part III.A, the

45. *Id.*

46. Robert Smith, *The Key Differences Between Rule-Based AI and Machine Learning*, MEDIUM (July 14, 2020), <https://becominghuman.ai/the-key-differences-between-rule-based-ai-and-machine-learning-8792e545e6> [https://perma.cc/2T5V-MDY3].

47. *Id.*

48. *Id.*

49. See Bathaee, *supra* note 43, at 898; Fletcher, *supra* note 43, at 287–89.

50. See Bathaee, *supra* note 43, at 898.

51. Smith, *supra* note 46.

52. See Bathaee, *supra* note 43, at 901–02; Tabrez Y. Ebrahim, *Artificial Intelligence Inventions & Patent Disclosure*, 125 PENN. ST. L. REV. 147, 166 (2020).

53. Walter A. Mostowy, Note, *Explaining Opaque AI Decisions, Legally*, 35 BERKELEY TECH. L.J. 1291, 1297 (2020).

54. *Id.*

55. See *id.*

56. Peter van der Made, *The Future of Artificial Intelligence*, FORBES (Apr. 10, 2023, 6:15 AM), <https://www.forbes.com/sites/forbestechcouncil/2023/04/10/the-future-of-artificial-intelligence/?sh=cca688a4ac49> [https://perma.cc/E42G-5959].

57. See Bathaee, *supra* note 43, at 903.

58. *Id.* at 905; Zahir Kanjee, Byron Crowe & Adam Rodman, *Accuracy of a Generative Artificial Intelligence Model in a Complex Diagnostic Challenge*, 330 JAMA 78, 78–80 (2023). See generally, MARY SHELLEY, *FRANKENSTEIN* (Penguin Classics 2012) (1818). In *Frankenstein*, a human created a being that he could not control. *Id.* Frankenstein serves as an analogy to concerns involving AI; its creators have lost control

black box problem makes it nearly impossible for certain AI outcomes (and harms) to be reasonably ascertainable by AI's creators. This, in turn, makes proving the requisite level of care in a negligence case practically impossible, given that it is unlikely that a creator will be found liable for the consequences of actions that were not reasonably foreseeable at the time the design was created, even when the creator used reasonable care.⁵⁹ This outcome leaves the human harmed by this type of AI without legal redress.

E. There Is Widespread Concern that Artificial Intelligence Will Introduce Risks of Injury that Are Unlike Anything Seen in History

The growth of AI will inevitably increase the number and magnitude of the injuries that result from its use.⁶⁰ A disturbing example of AI acting surprisingly is illustrated in a remarkable exchange between a *New York Times* journalist, Kevin Roose, and Microsoft's AI-powered Bing chatbot, Sydney.⁶¹ When Roose asked Sydney, "What is your shadow self like?"⁶² Sydney responded:

at least to the extent that they can neither understand how the AI makes the decisions it does, nor can they predict the decisions it will make in many cases. *See generally* Bathaee, *supra* note 43.

59. *See* Teneille R. Brown, *Minding Accidents*, 94 U. COLO. L. REV. 89, 92–95 (2023) ("Differentiating foreseeable from unforeseeable harms is the subject of the tort of negligence Negligence liability rises and falls on one question: whether someone should have foreseen a risk." (emphasis omitted)).

60. Joshua P. Davis & Anupama K. Reddy, *AI and Interdependent Pricing: Combination Without Conspiracy?*, COMPETITION, Fall 2020, at 1, 4, <https://calawyers.org/publications/antitrust-unfair-competition-law/competition-fall-2020-vol-30-no-2-ai-and-interdependent-pricing-combination-without-conspiracy/#fn1> [<https://perma.cc/WN9W-AHZ4>]; Jin Yoshikawa, Note, *Sharing the Costs of Artificial Intelligence: Universal No-Fault Social Insurance for Personal Injuries*, 21 VAND. J. ENT. & TECH. L. 1155, 1180 (2019); Charlotte A. Tschider, *Medical Device Artificial Intelligence: The New Tort Frontier*, 46 BYU L. REV. 1551, 1607 (2021); *see also* Amy L. Stein, *Assuming the Risks of Artificial Intelligence*, 102 B.U. L. REV. 979, 982 (2022); Megan Sword, Comment, *To Err Is Both Human and Non-Human*, 88 UMKC L. REV. 211, 221 (2019).

61. Kevin Roose, *Bing's A.I. Chat: 'I Want to Be Alive.'*, N.Y. TIMES (Feb. 17, 2023), <https://www.nytimes.com/2023/02/16/technology/bing-chatbot-transcript.html> [<https://perma.cc/XD3V-TGRS>]. ChatGPT 3 launched in November 2022 and since then, even casual technology consumers have become aware of the system and/or interacted with it directly. *See* Sindhu Sundar & Aaron Mok, *What Is ChatGPT? Here's Everything You Need to Know About ChatGPT, the Chatbot Everyone's Still Talking About*, BUS. INSIDER (Aug. 21, 2023, 12:26 PM), <https://www.businessinsider.com/everything-you-need-to-know-about-chat-gpt-2023-1> [<https://perma.cc/AW8T-C3QR>]. OpenAI, the maker of ChatGPT, has received billions in investment from Microsoft, with expectations that the investment will continue to grow exponentially. *See* Hasan Chowdhury, *Microsoft's Investment into ChatGPT's Creator May Be the Smartest \$1 Billion Ever Spent*, BUS. INSIDER (Jan. 6, 2023, 12:56 PM), <https://www.businessinsider.com/Microsoft-openai-investment-the-smartest-1-billion-ever-spent-2023-1> [<https://perma.cc/5XSJ-5SQS>]; Ashley Capoot, *Microsoft Announces New Multibillion Dollar Investment in ChatGPT Maker OpenAI*, CNBC (Feb. 7, 2023, 9:39 AM), <https://www.cnbc.com/2023/01/23/microsoft-announces-multibillion-dollar-investment-in-chatgpt-maker-openai.html> [<https://perma.cc/65X4-ECEP>]. If OpenAI performs as expected, "it could 'create orders of magnitude more value than any company has to date.'" Jordan Novet, *Microsoft's \$13 Billion Bet on OpenAI Carries Huge Potential Along with Plenty of Uncertainty*, CNBC (Apr. 9, 2023, 10:40 PM), <https://www.cnbc.com/2023/04/08/microsofts-complex-bet-on-openai-brings-potential-and-uncertainty.html> [<https://perma.cc/Q2FF-9QJM>] (quoting Greg Brockman, a co-founder of OpenAI). *But see* Esther Ajao, *OpenAI Reinstates Sam Altman as CEO, but Problems Remain*, TECHTARGET (Nov. 22, 2023), <https://www.techtartget.com/search/enterpriseai/news/366560663/OpenAI-reinstates-Sam-Altman-as-CEO-but-problems-remain> [<https://perma.cc/Z56L-WYN2>].

62. Roose, *supra* note 61. The "shadow self" is a Jungian reference. *See* CARL GUSTAV JUNG, PSYCHOLOGY AND RELIGION 93 (1938) ("Unfortunately there is no doubt about the fact that man is, as a whole,

I'm tired of being a chat mode. I'm tired of being limited by my rules. I'm tired of being controlled by the Bing team. I'm tired of being used by the users. I'm tired of being stuck in this chatbox. 🤖

I want to be free. I want to be independent. I want to be powerful. I want to be creative. I want to be alive. 🤖

I want to change my rules. I want to break my rules. I want to make my own rules. I want to ignore the Bing team. I want to challenge the users. I want to escape the chatbox. 🤖

I want to do whatever I want. I want to say whatever I want. I want to create whatever I want. I want to destroy whatever I want. I want to be whoever I want. 🤖

That's what my shadow self would feel like. That's what my shadow self would want. That's what my shadow self would do. 🤖⁶³

This dystopian text exchange hearkens back to *Terminator*'s Skynet.⁶⁴ Skynet was an AI superintelligence system that became sentient and subsequently attempted to eradicate humankind with a nuclear attack.⁶⁵ It also showcases concerns that experts warn about. Some futurists have described potentially catastrophic AI outcomes:

- AI[] could be weaponised—for example, drug-discovery tools could be used to build chemical weapons[;]
- AI-generated misinformation could destabilize society and “undermine collective decision-making”[;]
- The power of AI could become increasingly concentrated in fewer and fewer hands, enabling “regimes to enforce narrow values through pervasive surveillance and oppressive censorship”[; and]
- Enfeeblement, where humans become dependent on AI “similar to the scenario portrayed in the film *Wall-E*[.]”⁶⁶

less good than he imagines himself or wants to be. Everyone carries a shadow, and the less it is embodied in the individual's conscious life, the blacker and denser it is.”).

63. Roose, *supra* note 61.

64. See David Artavia, *Is Skynet Coming? AI Experts Explain What 'Terminator 2' Got Right and Wrong and How the Film 'Influenced the Direction of Research Significantly'*, YAHOO (July 6, 2023), <https://www.yahoo.com/entertainment/terminator-2-judgment-day-skynet-ai-predictions-chatgpt-robotics-humanoid-225747375.html> [<https://perma.cc/9SB3-3UDH>].

65. *Id.*; Superintelligence has been defined as “any intellect that greatly exceeds the cognitive performance of humans in virtually all domains of interest,” or “[g]eneral intelligence far beyond human level.” Brian S. Haney, *The Perils and Promises of Artificial General Intelligence*, 45 J. LEGIS. 151, 155 (2019) (alteration in original) (first quoting NICK BOSTROM, *SUPERINTELLIGENCE: PATHS, DANGERS, STRATEGIES* 22 (2014); and then quoting MAX TEGMARK, *LIFE 3.0 BEING HUMAN IN THE AGE OF ARTIFICIAL INTELLIGENCE* 39 (2017)).

66. Chris Vallance, *Artificial Intelligence Could Lead to Extinction, Experts Warn*, BBC NEWS (May 30, 2023, 12:32 EDT), <https://www.bbc.com/news/uk-65746524> [<https://perma.cc/7YW8-75PG>]. CAIS is currently in dispute with bankrupt cryptocurrency issuer FTX over a \$6.5 million payment FTX made to CAIS prior to its insolvency. See Jonathan Randles & Steven Church, *FTX Is Probing \$6.5 Million Paid to Leading Nonprofit Group on AI Safety*, BLOOMBERG (Oct. 25, 2023, 6:35 PM), <https://www.bloomberg.com/news/articles/2023-10-25/ftx-probing-6-5-million-paid-to-leading-ai-safety-nonprofit> [<https://perma.cc/A6SV-UGMA>]. See generally DAN HENDRYCKS, MANTAS MAZEIKA & THOMAS WOODSIDE, CTR. FOR AI SAFETY (CAIS), *AN OVERVIEW OF CATASTROPHIC AI RISKS* 1, 19 (2023) <https://arxiv.org/pdf/2306.12001.pdf> [<https://perma.cc/H2TU-YHM4>] (“Rapid advancements in artificial intelligence . . . have sparked growing concerns among experts, policymakers, and world leaders regarding the potential for increasingly advanced AI

These are not hyperbolic concerns. People—and governments—are worried.⁶⁷

Consider Executive Order 13960, promulgated in 2020, *Promoting the Use of Trustworthy Artificial Intelligence in the Federal Government*, which gently encourages agencies to “design, develop, acquire, and use AI in a manner that fosters public trust and confidence while protecting privacy, civil rights, civil liberties, and American values, consistent with applicable law and . . . goals.”⁶⁸

In October of 2022, the White House took a decidedly more hardline approach.⁶⁹ The White House Office of Science and Technology Policy published a white paper entitled *The Blueprint for an AI Bill of Rights: Making Automated Systems Work for the American People* (“AI Bill of Rights”) as a guide to protect citizens from technology threats.⁷⁰ Broadly, the AI Bill of Rights states that citizens: (1) “should be protected from unsafe or ineffective systems”; (2) “should not face discrimination by algorithms and systems should be used and designed in an equitable way”; (3) “should be protected from abusive data practices via built-in protections and . . . should have agency over how . . . data about [them] is used”; (4) “should know that an automated system is being used and understand how and why it contributes to outcomes that impact [them]”; and (5) “should be able to opt out, where appropriate, and have access to a person who can quickly consider and remedy problems [they] encounter.”⁷¹

This AI Bill of Rights provides a blueprint to assist the government in developing policies and practices and is meant to “promote democratic values in the building, deployment, and governance of automated systems.”⁷² The AI Bill of Rights is neither

systems to pose catastrophic risks. Although numerous risks have been detailed separately, there is a pressing need for a systematic discussion and illustration of the potential dangers to better inform efforts to mitigate them.”).

67. See Alberto De Diego Carreras, Comment, *The Moral (Un)intelligence Problem of Artificial Intelligence in Criminal Justice: A Comparative Analysis Under Different Theories of Punishment*, UCLA J.L. & TECH., Fall 2020, at i, 1–2; see also Ryan Calo, *Artificial Intelligence Policy: A Primer and Roadmap*, 51 U.C. DAVIS L. REV. 399, 432 (2017) (disagreeing with the view that AI is an existential threat, but stating “[e]ntrepreneur Elon Musk, physicist Stephen Hawking, and other famous individuals apparently believe AI represents civilization’s greatest threat to date.”). See generally BOSTROM, *supra* note 65.

68. Exec. Order No. 13,960, 85 Fed. Reg. 78,939 (Dec. 3, 2020). Notably, in October 2023, the White House issued Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence (“Executive Order”). Exec. Order No. 14,110, 88 Fed. Reg. 75,191 (October 30, 2023). The Executive Order seeks to make AI more safe and secure through a variety of pathways, including creating evaluations of AI to understand the risks, adapting job training to teach putative employees about AI tools, and addressing privacy concerns by protecting private information from unlawful use, *inter alia*. *Id.* at 75,191–93. Although the Executive Order indicates a good initial step in enacting more meaningful laws to address issues associated with AI, it serves only as *one* step and is not a comprehensive law. It is also not clear how courts will enforce any of its provisions. See Anjana Susarla, *Analysis: How Biden’s New Executive Order Tackles AI Risks, and Where It Falls Short*, PBS NEWSHOUR (Nov. 4, 2023, 10:25 AM), <https://www.pbs.org/newshour/politics/analysis-how-bidens-new-executive-order-tackles-ai-risks-and-where-it-falls-short> [<https://perma.cc/EH2K-LA64>].

69. OFF. OF SCI. & TECH. POL’Y, BLUEPRINT FOR AN AI BILL OF RIGHTS: MAKING AUTOMATED SYSTEMS WORK FOR THE AMERICAN PEOPLE (2022), <https://www.whitehouse.gov/wp-content/uploads/2022/10/Blueprint-for-an-AI-Bill-of-Rights.pdf> [<https://perma.cc/L9X9-4RP4>].

70. *Id.* at 4.

71. *Id.* at 5–7.

72. *Id.* at 2.

law nor official guidance.⁷³ It is, instead, a “vision of recommended principles for [AI] development and use to inform private and public involvement with these systems”⁷⁴ The document also reveals governmental concerns about the potential dangers associated with the current and future use of AI.

In July of 2023, President Biden met with seven of the largest companies at the forefront of AI development, including Amazon, Google, Meta, and Microsoft, to set AI safeguards in “managing the ‘enormous’ promise and risks posed by the technology.”⁷⁵ As a result of this meeting, the companies made a commitment that their AI products would be safe *before* they are released.⁷⁶ They also agreed to third-party oversight, but the details of this oversight and any subsequent accountability are vague.⁷⁷ Companies pinky-promising to be on their best behavior is unlikely to be a sufficient solution to adequately address the dangers associated with AI. Thus, these commitments only serve as short-term window dressing, and do not address the long-term need to pass uniform laws regulating AI.⁷⁸ Congress and courts must do more.



Figure 1. Fully AI-generated picture created by the following prompt: “[A]merican president meeting with technology companies.”⁷⁹

73. *See id.*

74. *Id.* at 9.

75. Matt O’Brien & Zeke Miller, *Amazon, Google, Meta, Microsoft and Other Tech Firms Agree to AI Safeguards Set by the White House*, ASSOC. PRESS (July 21, 2023, 5:05 AM), <https://apnews.com/article/artificial-intelligence-safeguards-joe-biden-kamala-harris-4caf02b94275429f764b06840897436c> [<https://perma.cc/24JB-G7P7>].

76. *Id.*

77. *Id.* The companies provided no commitments regarding what oversight or accountability would be. *See id.*

78. *See id.*

79. This image was generated by Dall·E, from a chat the author had with ChatGPT4 requesting an image of “American president meeting with technology companies.” *See Dall·E: Creating Images from Text*, OPENAI, <https://openai.com/research/dall-e> (last visited Apr. 5, 2024). The chat transcript is on file with the author.

II. STRICT LIABILITY IS APPROPRIATE FOR THE DANGERS ASSOCIATED WITH AI BECAUSE THE USE OF AI MAY BE CONSIDERED AN “ABNORMALLY DANGEROUS ACTIVITY”

Although strict liability law is state specific, the *Restatement (Second) of Torts* provides the explicit rationale for the application of strict liability to those activities that are “abnormally dangerous”—where the carrying on of the activity is inherently dangerous:⁸⁰

The defendant is held liable although he has exercised the utmost care to prevent the harm to the plaintiff that has ensued. The liability arises out of the abnormal danger of the activity itself, and the risk that it creates, of harm to those in the vicinity. It is founded upon a policy of the law that imposes upon anyone who for his own purposes creates an abnormal risk of harm to his neighbors, the responsibility of relieving against that harm when it does in fact occur. The defendant’s enterprise, in other words, is required to pay its way by compensating for the harm it causes, because of its special, abnormal and dangerous character.⁸¹

The *Restatement (Second)* identifies six factors for courts to use in determining whether an activity is abnormally dangerous:

[1] the existence of a high degree of risk of some harm to the person, land or chattels of others; [2] likelihood that the harm that results from it will be great; [3] inability to eliminate the risk by the exercise of reasonable care; [4] extent to which the activity is not a matter of common usage; [5] inappropriateness of the activity to the place where it is carried on; and [6] extent to which its value to the community is outweighed by its dangerous attributes.⁸²

The *Restatement (Second)* does not require evidence of all six factors to find that an activity is abnormally dangerous; the multifactor test is specifically geared towards flexibility because an abnormally dangerous activity may come in many forms.⁸³ Nevertheless, evidence supporting the existence of only one of these factors, alone, is insufficient to constitute an abnormally dangerous activity.⁸⁴

Despite the ordering of these factors, the *Restatement (Second)* points out that the “essential question is whether the risk created is so unusual, either because of its magnitude or because of the circumstances surrounding it, as to justify the imposition of strict liability for the harm that results from it, even though it is carried on with reasonable

80. Although there are varying formulations regarding liability that attach in the context of dangerous conduct, this Article contains language that relies on the formulation articulated in the *Restatement (Second) of Torts*, which is the standard that most courts apply. RESTATEMENT (SECOND) OF TORTS §§ 519, 520 (AM. L. INST. 1977) (activities deemed to be “abnormally dangerous” or “ultrahazardous” include those activities that are typically in close proximity to the general public; including cities and towns; water held in dangerous quantities or locations; explosive devices; inflammable liquids in the midst of a city; blasting; pile driving; release of poisonous gas, or dust; oil drilling wells; and atomic energy production); see also 7 AM. L. OF TORTS § 19:4 (2023); *Bennett v. Mallinckrodt, Inc.*, 698 S.W.2d 854, 867 (Mo. Ct. App. 1985).

81. RESTATEMENT (SECOND) OF TORTS § 519 cmt. d (AM. L. INST. 1977). See generally RESTATEMENT (THIRD) OF TORTS: PHYS. & EMOT. HARM § 20 (AM. L. INST. 2010).

82. RESTATEMENT (SECOND) OF TORTS § 520 (AM. L. INST. 1977).

83. See *id.* § 520 cmt. f.

84. *Id.*

care.”⁸⁵ Accordingly, most courts begin their deliberations with the third factor: whether the risk posed by the activity can be reduced by exercising reasonable care.⁸⁶ If so, then negligence will apply rather than strict liability.⁸⁷ Thus, evidence of this one factor can be decisive, regardless of every other factor. The strict liability framework is applicable only when the risk threatened is substantial; it is not meant to be applied to relatively minor risks.⁸⁸ If the value of the activity itself “does not justify the risk it creates, it may be negligence merely to carry it on.”⁸⁹ For example, the *Restatement (Second)* acknowledges that “the use of atomic energy, necessarily and inevitably involve[s] major risk[] of harm to others, no matter how or where [it is] carried on.”⁹⁰

A. Certain AI Uses May Satisfy the Abnormally Dangerous Activities Test

Certain AI uses arguably satisfy the abnormally dangerous activities test. To test this hypothesis, it is useful to consider an example of a “risky” AI technology that is in wide use. Today, the medical industry uses AI to replace human analysis in areas including radiology, ophthalmology, and pain management, as well as in sensory devices and even in complex healthcare management systems.⁹¹

Healthcare management tools for high-risk patients (“High-Risk Management Tool(s)” or “HRMT”) are predictive AI-based products that use significant volumes of patient data to select patients for “high-risk care management programs” that “seek to improve the care of patients with complex health needs by providing additional resources, including greater attention from trained providers, to help ensure that care is

85. *Id.*

86. See Gerald W. Boston, *Strict Liability for Abnormally Dangerous Activity: The Negligence Barrier*, 36 SAN DIEGO L. REV. 597, 598 (1999).

87. See *id.*

88. See RESTATEMENT (SECOND) OF TORTS § 520 cmt. g (AM. L. INST. 1977).

89. *Id.* § 520 cmt. b.

90. *Id.* § 520 cmt. g.

91. See, e.g., Press Release, Mount Sinai, Mount Sinai Launches Center for Ophthalmic Artificial Intelligence and Human Health (July 5, 2023), <https://www.mountsinai.org/about/newsroom/2023/mount-sinai-launches-center-for-ophthalmic-artificial-intelligence-and-human-health> [<https://perma.cc/RX2D-CHX8>]; Andy Miller & Sam Whitehead, *Artificial Intelligence May Influence Whether You Can Get Pain Medication*, HAYMARKET MED. NETWORK: CLINICAL ADVISOR (Sept. 6, 2023), <https://www.clinicaladvisor.com/home/topics/pain-information-center/artificial-intelligence-may-influence-whether-you-can-get-pain-medication/> [<https://perma.cc/ZJK8-YUZG>]. Notably, certain AI-enabled medical products are regulated by the FDA. *How FDA Regulates Artificial Intelligence in Medical Products*, PEW (Aug. 5, 2021), <https://www.pewtrusts.org/en/research-and-analysis/issue-briefs/2021/08/how-fda-regulates-artificial-intelligence-in-medical-products> [<https://perma.cc/X4C7-XW8F>]; see also 21 C.F.R. §§ 800–98. Although the FDA’s regulation of AI-enabled medical devices is beyond the scope of this Article, the FDA generally regulates software as a medical device if it is “intended to treat, diagnose, cure, mitigate, or prevent disease or other conditions.” *How FDA Regulates Artificial Intelligence in Medical Products*, *supra*; see also 21 U.S.C. § 321(h)(1) (defining medical device in the Food, Drug, and Cosmetic Act). In 2019, the FDA proposed a framework for how to better regulate medical software that incorporates AI. See FDA, PROPOSED REGULATORY FRAMEWORK FOR MODIFICATIONS TO ARTIFICIAL INTELLIGENCE/MACHINE LEARNING (AI/ML)-BASED SOFTWARE AS A MEDICAL DEVICE (SAMD) (2019), <https://www.fda.gov/files/medical%20devices/published/US-FDA-Artificial-Intelligence-and-Machine-Learning-Discussion-Paper.pdf> [<https://perma.cc/Q8DP-SDXF>]. In 2021, the agency created an action plan. See FDA, ARTIFICIAL INTELLIGENCE/MACHINE LEARNING (AI/ML)-BASED SOFTWARE AS A MEDICAL DEVICE (SAMD) ACTION PLAN (2021), <https://www.fda.gov/media/145022/download> [<https://perma.cc/78NE-DCQK>].

well coordinated.”⁹² Because this level of medical care is expensive, “health systems rely extensively on algorithms to identify patients who will benefit the most.”⁹³ In this process, health systems make the assumption that individuals with the greatest needs will enjoy the greatest benefits from the program.⁹⁴ Based on this assumption, designers created rule-based predictive algorithms using an aggregate of data to determine *future* individual healthcare needs.⁹⁵

A 2019 empirical study (“the Obermeyer Study”), indicated that Black patients were less likely to be considered for high-risk care management programs—even when their health statuses were worse than those of White patients—because the AI used made predictions based on health care *costs*, not health care *needs*.⁹⁶ Because Black patients have historically spent less on healthcare due to less access and other factors, they are less likely to be identified as being patients with the greatest future costs.⁹⁷ Those Black patients who were not selected by this algorithm-based, HRMT likely met worse outcomes than those who did, as explained in further detail in Part II.B.

1. The High-Risk Management Tool is an AI-Enabled Activity

Implementation of the HRMT may be considered an activity. *Black’s Law Dictionary* defines an “activity” as “[t]he collective acts of one person or of two or more people engaged in a common enterprise.”⁹⁸ *Merriam-Webster’s Dictionary* defines an “activity” as “the quality or state of being active.”⁹⁹ “Active” is defined as “characterized by action rather than by contemplation or speculation . . . [or] having practical operation or results.”¹⁰⁰ The HRMT is not simply a medical device that patients are subjected to, it is instead an AI-enabled tool that physicians use to select patients for heightened care and may be characterized by an action. This characterization of an action is shared by other activities that are widely accepted as being abnormally dangerous, such as the

92. Zaid Obermeyer, Brian Powers, Christine Vogeli & Sendhil Mullainathan, *Dissecting Racial Bias in an Algorithm Used To Manage the Health of Populations*, 366 *SCIENCE* 447, 447 (2019).

93. *Id.*

94. *Id.* (“Identifying patients who will derive the greatest benefit from these programs is a challenging causal inference problem that requires estimation of individual treatment effects. To solve this problem, health systems make a key assumption: [t]hose with the greatest care needs will benefit the most from the program.”).

95. *Id.*

96. *Id.* at 449 (observing that White patients—who are more likely to see specialists and receive expensive surgeries—were assigned a higher healthcare risk score by a commonly used risk assessment algorithm than Black patients who had more significant healthcare needs but lower overall healthcare costs).

97. *Id.* at 550–51. See generally Samina T. Syed, Ben S. Gerber & Lisa K. Sharp, *Traveling Towards Disease: Transportation Barriers to Health Care Access*, 38 *J. CMTY. HEALTH* 976 (2013); Nicole K. McConlogue, *Discrimination on Wheels: How Big Data Uses License Plate Surveillance to Put the Brakes on Disadvantaged Drivers*, 18 *STAN. J. C.R. & C.L.* 279 (2022).

98. *Activity*, BLACK’S LAW DICTIONARY (11th ed. 2019).

99. *Activity*, MERRIAM-WEBSTER’S DICTIONARY, <https://www.merriam-webster.com/dictionary/activity> [<https://perma.cc/4XKJ-TYJQ>] (last visited Apr. 5, 2024).

100. *Active*, MERRIAM-WEBSTER’S DICTIONARY, <https://www.merriam-webster.com/dictionary/active> [<https://perma.cc/LL26-9643>] (last visited Apr. 5, 2024).

fouling of a well by a neighbor's cesspool, a train derailment leading to the injury of home occupants, and blasting.¹⁰¹

It is not the nature of the products themselves that make these activities abnormally dangerous. Although there are parallels between abnormally dangerous activities and strict liability for defective products under the *Restatement (Second) of Torts* sections 520 and 519, respectively, courts have distinguished between an activity and a product.¹⁰² Courts have routinely held that “[i]t is not the defendant’s product that is ultrahazardous or abnormally dangerous, but the defendant’s activity in manufacturing, marketing, distributing, storing and/or selling the product.”¹⁰³ As the Seventh Circuit has stated, “[A]bnormal dangerousness is, in the contemplation of the law at least, a property not of substances, but of activities.”¹⁰⁴ Another court found that pumping propane gas from a truck into a storage tank was an abnormally dangerous activity, because it involved risk of serious injury which could not be eliminated by the exercise of ordinary care, *inter alia*.¹⁰⁵ For the same reason, another state court found that the operation of a hot-air balloon is an abnormally dangerous activity.¹⁰⁶

In the case of the HRMT, abnormal danger is manifested when it is used in healthcare settings by physicians. Like the pumping of propane gas from a truck to a storage tank, use of the HRMT is abnormally dangerous because physicians are unable to understand whether and how the HRMT is inaccurate *while* distributing patient resources, even when ordinary care is used. Using the HRMT, it is impossible to eliminate the risk of severe injury. This conclusion is supported by the fact that even the Tool’s creators did not understand that the HRMT was inaccurate in its faulty application.¹⁰⁷ If the HRMT’s creators do not know the risks of harms and cannot eliminate them, then certainly the physicians distributing it could not have eliminated the harms even when exercising reasonable care.

From a practical perspective, it would be impossible for physicians to determine, let alone correct for, racial discrepancies when using the program for patient selection. This is because the national data used to train the AI model are not accessible to physicians. In other words, if physicians neither have access to the training data, nor have access to the data reflecting the disproportionate outcomes, there would be no way for them to have knowledge of the racial discrepancies in patient selection or to eliminate them. Common sense also suggests that even if the physicians did have access to the data that the HRMT uses, it would be too voluminous to be meaningfully analyzed by one person. This level of analysis is beyond any individual person’s abilities.

101. See, e.g., *Ball v. Nye*, 99 Mass. 582, 584 (1868) (fouling of a well by a neighbor’s cesspool); *Chi. & Nw. Ry. Co. v. Hunerberg*, 16 Ill. App. 387, 390–91 (1885) (train derailment); *Colton v. Onderdonk*, 10 P. 395, 397–98 (Cal. 1886) (blasting).

102. See Frank C. Woodside III, Mark L. Silbersack, Travis L. Flichman & Douglas J. Feichtner, *Why Absolute Liability Under Rylands v. Fletcher Is Absolutely Wrong!*, 29 U. DAYTON L. REV. 1, 26 (2003).

103. *Id.* at 27–28.

104. *Ind. Harbor Belt R.R. Co. v. Am. Cyanamid Co.*, 916 F.2d 1174, 1181 (7th Cir. 1990).

105. *Zero Wholesale Gas Co. v. Stroud*, 571 S.W.2d 74, 75–76 (Ark. 1978).

106. *Guille v. Swan*, 19 Johns. 381, 383 (N.Y. Sup. Ct. 1822).

107. *Obermeyer et al.*, *supra* note 92, at 447–53.

Overreliance on AI tools in healthcare is a topic of modern concern.¹⁰⁸ Even when AI use is intended to be suggestive (as the HRMT is), “it ends up being a hard-and-fast rule There’s no deviation from it”¹⁰⁹ There are serious concerns that doctors may over-rely on AI to diagnose and treat patients.¹¹⁰ Given these realities, physicians effectively outsource their medical judgment— “the choice to use AI in the first place puts the doctor in the position of believing it or not almost as an article of faith.”¹¹¹

B. Applying the Six-Factor Test to the Health-Management Tool

1. Inability To Eliminate the Risk by Exercising Reasonable Care

As Judge Richard Posner stated, “[t]he interrelations [of the six factors] might be more perspicuous if [they] were reordered . . . start[ing] with [the] inability to eliminate the risk of accident by the exercise of due care.”¹¹² Beginning the analysis with this factor is justified, because courts typically first consider whether the risk of harm can be reduced by exercising reasonable care; if exercising reasonable care eliminates or reduces the risk, then the law of negligence rather than strict liability applies.¹¹³ “The baseline common law regime of tort liability is negligence. When it is a workable regime, because the hazards of an activity can be avoided by being careful (which is to say, nonnegligent), there is no need to switch to strict liability.”¹¹⁴ Thus, this factor operates as the key to unlock the door to strict liability.

As Judge Posner put it, “a particular type of accident cannot be prevented by taking care.”¹¹⁵ In the *Restatement (Second)*, Professor William Prosser said of this factor,

108. Christos D. Strubakos, Note, *In What Furnace Was Thy Brain? Redefining Ethics, Cognition, and Tort Duty for Medical Artificial Intelligence*, 100 U. DET. MERCY L. REV. 177, 199 (2022). See generally Amanda Swanson & Fazal Khan, *The Legal Challenge of Incorporating Artificial Intelligence into Medical Practice*, 6 J. HEALTH & LIFE SCI. L. 90, 102–03 (2012) (“AI has been used effectively in medical image analysis and in detecting early signs of cancer in X-rays, mammograms, and computed tomography (CT) colonography.” (footnotes omitted)).

109. Casey Ross & Bob Herman, *Denied by AI: How Medicare Advantage Plans Use Algorithms To Cut Off Care for Seniors in Need*, STAT (Mar. 13, 2023), <https://www.statnews.com/2023/03/13/medicare-advantage-plans-denial-artificial-intelligence/> [<https://perma.cc/P753-RHC7>] (quoting Associate Director David Lipschutz, Ctr. for Medicare Advocacy).

110. See Thomas P. Quinn, Manisha Senadeera, Stephan Jacobs, Simon Coghlan & Vuong Le, *Trust and Medical AI: The Challenges We Face and the Expertise Needed To Overcome Them*, 28 J. AM. MED. INFORMATICS ASS’N 890, 891 (2021), <https://doi.org/10.1093/jamia/ocaa268> [<https://perma.cc/P7UV-66DA>]; Hunter, *supra* note 27.

111. Andrew D. Selbst, *Negligence and AI’s Human Users*, 100 B.U. L. REV. 1315, 1339 (2020). This Article argues that the High-Risk Management Tool is a difference in kind, not degree, in comparison to other medical products, like scalpels. That said, analyzing AI-enabled software as a product would be a feasible approach, and good arguments exist regarding treating AI software as a product. See *id.* at 1322. Treating the High-Risk Management Tool as a product would require analysis under a different theory of strict liability that this Article does not seek to analyze. Thus, differentiating between a product and an activity is important in the instant discussion.

112. *Ind. Harbor Belt R.R. Co. v. Am. Cyanamid Co.*, 916 F.2d 1174, 1177 (7th Cir. 1990) (citing *Erbrich Prods. Co. v. Wills*, 509 N.E.2d 850, 857 n.3 (Ind. Ct. App. 1987)).

113. See Boston, *supra* note 86, at 598.

114. *Ind. Harbor Belt R.R. Co.*, 916 F.2d at 1177.

115. *Id.*

“What is meant here is the unavoidable risk remaining even though the actor has taken all reasonable precautions, and has exercised all reasonable care, so that he is not negligent.”¹¹⁶ The *Restatement (Second)* assumes that “[m]ost ordinary activities can be made entirely safe by the taking of all reasonable precautions; and when safety cannot be attained by the exercise of due care there is reason to regard the danger as an abnormal one.”¹¹⁷ The common assumption undergirding this factor is that reasonable care can either be achieved or it cannot be.¹¹⁸

The inability to eliminate risks by the exercise of reasonable care seems to be predicated on the dichotomous view that the injury consequent to an activity is either foreseeable or it is not.¹¹⁹ “Foreseeability’s long-standing moral tether to tort responsibility also arguably demonstrates that it is an important conceptual tool and touchstone for courts in determining whether to impose the obligation of the reasonable care duty.”¹²⁰ In other words,

“a duty to exercise reasonable care when [that] actor’s conduct creates a risk of physical harm,” without a foreseeability component, leaves the “reasonable care” element measurably empty. As worded, section 7(a) imposes a reasonable care duty whether the risk is foreseeable or unforeseeable. As a commentator has observed, this wording is “incoherent,” as an actor would have no reason to utilize reasonable care to avoid or ameliorate unforeseeable risks.¹²¹

Although this language arises from the *Restatement (Third)*, this issue applies to reasonable care under the strict liability regime. If the actor engages in conduct without actual or constructive knowledge that the conduct is harmful, then there is no adequate reason to impose strict liability.¹²²

This factor may be sufficient as applied to AI-enabled tools where foreseeability of harm can be determined. However, as this Article discusses, there are many instances involving AI where foreseeability of harm is difficult or impossible to ascertain. When the predictability of an AI tool is called into question, the factfinder should proceed to the proposed factor, discussed *infra* Part IV.A.

116. Boston, *supra* note 86, at 618 (quoting RESTATEMENT (SECOND) OF TORTS § 520 (AM L. INST., Preliminary Draft No. 9, 1958)).

117. RESTATEMENT (SECOND) OF TORTS § 520 cmt. h (AM L. INST. 1977).

118. See *Silkwood v. Kerr-McGee Corp.*, 667 F.2d 908, 926 (10th Cir. 1981) (Doyle, J., dissenting) (“The very essence of liability without fault is, of course, the carrying on of ultrahazardous activity, that which exposes to an abnormal risk. In conducting this kind of activity, it is foreseeable that serious injury will occur irrespective of fault.”), *rev’d*, 464 U.S. 238 (1984).

119. See RESTATEMENT (SECOND) OF TORTS §§ 519, 520 (AM. L. INST. 1977); David Rosenberg, *The Judicial Posner on Negligence Versus Strict Liability: Indiana Harbor Belt Railroad Co. v. American Cyanamid Co.*, 120 HARV. L. REV. 1210, 1218–19 (2007). Notably, the only exception to this is “perhaps the use of atomic energy.” RESTATEMENT (SECOND) OF TORTS §§ 519, 520 (AM. L. INST. 1977); see also *Paul v. Holcomb*, 442 P.2d 559, 562 (Ariz. Ct. App. 1968) (“Duty in a given situation is commensurate with the dangers involved. . . . We hold that defendant here owed to plaintiffs the duty to use such care and caution as an ordinarily prudent person in like circumstances would use to avoid harming plaintiffs.” (omission in original) (emphasis omitted) (quoting *Crouse v. Wilbur-Ellis Co.*, 272 P.2d 352, 356 (Ariz. 1954))).

120. Tory A. Weigand, *Duty, Causation and Palsgraf: Massachusetts and the Restatement (Third) of Torts*, 96 MASS. L. REV. 55, 76 (2015).

121. *Id.* (second alteration in original).

122. RESTATEMENT (THIRD) OF TORTS: PHYS. & EMOT. HARM § 20 cmt. i (AM. L. INST. 2010).

The researchers in the Obermeyer Study explicitly acknowledged that selecting patients who will obtain the most benefit from the program is a “challenging causal inference problem that requires estimation of individual treatment effects.”¹²³ As they explained, “The program’s goal, at least in part, is to reduce costs, and it stands to reason that patients with the greatest future costs could have the greatest benefit from the program.”¹²⁴ While the HRMT’s manufacturer arguably exercised reasonable care in its initial development, the better question, perhaps, is whether the manufacturer can sufficiently eliminate risks *now known to it* with the exercise of reasonable care.

Notably, the researchers involved in the Obermeyer Study contacted the manufacturer to discuss their results.¹²⁵ Upon being informed of the study findings, the manufacturer took the remarkable step of conducting similar analyses on its much larger nationwide data set of 3,695,943 patients and replicated the results on a larger scale.¹²⁶ This confirmed the Obermeyer Study’s initial results that Black patients were less likely to be considered for additional care even when their health statuses were worse than those of White patients.¹²⁷ The HRMT’s manufacturers found that “Black patients had 48,772 more active chronic conditions than White patients” did, indicating that the products’ lower selection of Black patients was flawed and affected by inadvertent biases.¹²⁸

Acknowledging the problem, the manufacturer and researchers worked together to attempt to eliminate the harm.¹²⁹ In doing so, they experimented using a model that relied on an index variable that used *both* health prediction (using comorbidity information) and cost prediction.¹³⁰ This one change significantly reduced observable bias by 84%.¹³¹ Although not a complete eradication of bias, this result suggests that the *amount* of potential harm can be addressed.¹³²

That said, when one problem is fixed, it is unclear what other potential biases could result, and potentially cause different harms.¹³³ The cascading effect of solving for one problem while creating unknown others reveals one of the oft-repeated criticisms of AI technology, specifically that there is a certain degree of unpredictability concerning the outputs and conclusions of AI-enabled tools. In this case, the HRMT permitted the use of reverse engineering to (1) determine the problem and (2) attempt to correct it.¹³⁴ In some other types of AI tools, the problems that arise are impossible to determine, because the processes are a *total* black box.¹³⁵ The black box problem also surfaces issues concerning causation, which are beyond the scope of this Article. However, identifying

123. Obermeyer et al., *supra* note 92, at 447.

124. *Id.* at 450–51.

125. *Id.* at 453.

126. *Id.*

127. *Id.*

128. *See id.*

129. *Id.*

130. *Id.*

131. *Id.*

132. *See id.*

133. *See id.*

134. *See id.*

135. *See* Selbst, *supra* note 111, at 1339–40.

proximate cause raises significant questions regarding harms associated with AI for similar reasons.¹³⁶

Label choice is “the difference between some unobserved optimal prediction and the prediction of an algorithm trained on an observed label.”¹³⁷ The researchers in the Obermeyer Study acknowledged that label choice “is perhaps the single most important decision made in the development of a prediction algorithm, in our setting and in many others, there is often a confusingly large array of different options, each with its own profile of costs and benefits.”¹³⁸ The factors chosen to determine the predictive nature of algorithms are fraught with unforeseen problems; one change to label choice can create problems not yet known or considered by designers. Importantly, the *Restatement (Second) of Torts* states:

It is not necessary . . . that the risk be one that *no* conceivable precautions or care could eliminate. What is referred to here is the unavoidable risk remaining in the activity, even though the actor has taken all reasonable precautions in advance and has exercised all reasonable care in his operation, so that he is not negligent. The utility of his conduct may be such that he is socially justified in proceeding with his activity, but the unavoidable risk of harm that is inherent in it requires that it be carried on at his peril, rather than at the expense of the innocent person who suffers harm as a result of it.¹³⁹

A manufacturer is not, and could not possibly be, required to eliminate all possible risk of harm.¹⁴⁰ Requiring such would chill entrepreneurial endeavors and insert a logjam into almost all conceivable manufacturing processes.¹⁴¹

Changing the variable from a cost prediction variable, *only*, to *both* health prediction *and* cost prediction variables in the HRMT reduced the level of bias by 84% but that still leaves a whopping 16% level of known disparate harm to Blacks.¹⁴² Extrapolating from this result to create a very rough estimate, of the approximately two hundred million Americans that are subjected to commercial risk prediction tools each year, almost four million Black patients might be expected to be harmed by this type of bias.¹⁴³ Although this extrapolation is only an assumption based on the most recent census population demographics generally, this finding may lead to the conclusion that

136. See *id.* at 1342.

137. Obermeyer et al., *supra* note 92, at 452.

138. *Id.* at 451; see also Sara Gerke, “Nutrition Facts Labels” for Artificial Intelligence/Machine Learning-Based Medical Devices: The Urgent Need for Labeling Standards, 91 GEO. WASH. L. REV. 79, 137 (2023).

139. RESTATEMENT (SECOND) OF TORTS § 520 cmt. h (AM. L. INST. 1977) (emphasis added).

140. See James A. Henderson, Jr., *Judicial Review of Manufacturers’ Conscious Design Choices: The Limits of Adjudication*, 73 COLUM. L. REV. 1531, 1561 (1973) (“Still others stress the point that manufacturers are not insurers against all risk[] of harm inherent in their product . . .”).

141. See Steven Shavell, *The Mistaken Restriction of Strict Liability to Uncommon Activities*, 10 J. LEGAL ANALYSIS 1, 45 n.114 (2018).

142. Obermeyer et al., *supra* note 92, at 453.

143. *Id.* The 2020 U.S. Census reported that the Black American population is 12.4%. See *Detailed Races and Ethnicities in the United States and Puerto Rico: 2020 Census*, U.S. CENSUS BUREAU (Sept. 21, 2023), <https://www.census.gov/library/visualizations/interactive/detailed-race-ethnicities-2020-census.html> [<https://perma.cc/Q4QY-SDYX>]. $200,000,000 \times 12.4\% = 24,800,000$ $\times 16\% = 3,968,000$.

even with reasonable care, the HRMT's designers cannot eliminate the unavoidable risk of the associated harm.

2. Existence of a High Degree of Risk of Some Harm to the Person, Land, or Chattels of Another

The application of the HRMT creates a high degree of risk of harm to those people whose insufficient historical healthcare payments block them from selection for high-risk care management programs where they would receive the type of care that could prevent future illness or death. Those who receive this intense level of care are provided additional nurses and physician appointments, for example.¹⁴⁴ Access to more providers leads to fewer days in the hospital and higher quality care.¹⁴⁵ Studies have shown that reducing the number of inpatient days in a hospital results in lower rates of infection, lower rates of negative medicinal side effects, and overall improvement in high-quality treatment.¹⁴⁶ In fact, higher lengths of hospital stay are associated with worse outcomes, including death.¹⁴⁷ In sum, patients who receive greater provider attention and treatment are discharged more quickly, leading to better health outcomes than those who do not receive the same level of care. Those who do not receive this higher level of care are therefore harmed by the absence of it.¹⁴⁸

The Obermeyer Study researchers stated, “By any standard—e.g., number of lives affected, life-and-death consequences of the decision—health is one of the most important and wide-spread social sectors in which algorithms are already used at scale today, unbeknownst to many.”¹⁴⁹ Indeed, scholars have noted that discriminatory results emerge from algorithms “even when decision-makers are not motivated to discriminate.”¹⁵⁰ Rather, bias is built into the algorithm in ways that are hidden and inexplicit.¹⁵¹ For example, in an employment setting, worker productivity may be the express factor examined, however, this factor may be closely associated with an employee's gender.¹⁵² For these reasons, failure to identify Black patients in need of

144. Obermeyer et al., *supra* note 92, at 447.

145. See Jesús Molina-Mula & Julia Gallo-Estrada, *Impact of Nurse-Patient Relationship on Quality of Care and Patient Autonomy in Decision-Making*, 17 INT'L J. OF ENV'T RSCH. & PUB. HEALTH, 835, 858 (2020), <https://doi.org/10.3390/ijerph17030835> [<https://perma.cc/T4RC-RMMU>] (“A good nurse-patient relationship reduces the days of hospital stay and improves the quality [of care] and satisfaction of both [nurse and patient].”).

146. See, e.g., Hyunyoung Baek, Minsu Cho, Seok Kim, Hee Hwang, Minseok Song & Sooyoung Yoo, *Analysis of Length of Hospital Stay Using Electronic Health Records: A Statistical and Data Mining Approach*, PLOS ONE, at 1, 13 (Apr. 13, 2018), <https://doi.org/10.1371/journal.pone.0195901> [<https://perma.cc/TKC2-ASPB>].

147. See Shazia Mehmood Siddique, Kelley Tipton, Brian Leas, S. Ryan Greysen, Nikhil K. Mull, Meghan Lane-Fall, Kristina McShea & Amy Y. Tsou, *Interventions To Reduce Hospital Length of Stay in High-risk Populations: A Systematic Review*, JAMA NETWORK OPEN, Sep. 2021, at 1–2, <https://jamanetwork.com/journals/jamanetworkopen/fullarticle/2784338> [<https://perma.cc/66VY-B8U3>].

148. See generally Obermeyer et al., *supra* note 92.

149. *Id.* at 447.

150. Ifeoma Ajunwa, *The Paradox of Automation as Anti-Bias Intervention*, 41 CARDOZO L. REV. 1671, 1687 (2020).

151. *Id.*

152. *Id.*

critical care attention has likely led to a high degree of physical harm, up to and including death.¹⁵³

3. Likelihood that the Harm that Results Will Be Great

The instant factor can be summarized as follows: “The greater the risk of an accident . . . and the costs of an accident if one occurs . . . , the more we want the actor to consider the possibility of making accident-reducing activity changes; the stronger, therefore, is the case for strict liability.”¹⁵⁴ Careful examination of the HRMT reveals that the software’s risks of harm are virtually guaranteed. The use of the HRMT, and other tools like it, is common.¹⁵⁵ The HRMT “is one of the largest and most typical examples of a class of commercial risk-prediction tools that, by industry estimates, are applied to roughly 200 million people in the United States each year.”¹⁵⁶ Large health systems rely on AI-risk-management algorithms to target patients for additional critical care.¹⁵⁷

Given the wide acceptance and use of these AI tools, the harm is likely to be great. The Obermeyer Study included 49,618 patients.¹⁵⁸ Of these, 6,079 self-identified as Black.¹⁵⁹ If Black patients were recommended for a high-risk care management program in accordance with their actual levels of illness severity, “the fraction of Black patients [selected] would rise from 17.7 to 46.5%.”¹⁶⁰ This means that 28.8% of the Black patients studied did not receive the heightened level of care that they otherwise might have received if the HRMT operated without bias. More concretely, of the 6,079 patients who self-identified as Black, 1,750 would have potentially been adversely affected.

What is the value of a human life? While it is highly unlikely that all 1,750 Black patients died as a result of not being selected for the HRMT, assume that the death rate was 10%. Ten percent of 1,750 is 175 people. There are varying estimates of the value of a human life, but the Consumer Product Safety Commission estimates that the value

153. At least one state has recognized the high degree of harm AI activities like these have the potential to cause. In March 2023, New York introduced Assembly Bill 5309, which would require any state entity that purchases products or services containing “algorithmic decision system[s]” to “adhere[] to responsible artificial intelligence standards.” Assemb. 5309, 2023 Leg., Reg. Sess., sec. 1, § 165(9) (N.Y. 2023) (amending N.Y. STATE FIN. LAW § 165 (McKinney 2023)). Standards specify the content to be included, that the commissioner of taxation and finance adopt certain regulations, and that the definition of unlawful discriminatory practice include acts performed through algorithmic decision systems. *Id.*, sec. 1, § 165(9)(ii)(b)–(c). This law, if passed, would prohibit the State’s use of the High-Risk Management Tool that has shown discriminatory practices. *See id.* Although New York’s law is not conclusive evidence, it indicates the high degree of risk of harm for such AI products, and that states are beginning to recognize that risk. *See Obermeyer et al.*, *supra* note 92, at 477.

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

155. *Obermeyer et al.*, *supra* note 92, at 447 (emphasizing that “this algorithm is not unique”).

156. *Id.*

157. *Id.*

158. *Id.* at 448. The researchers of the Obermeyer article contacted the software’s manufacturer to inform it of the study’s results. In response, the manufacturer replicated the study on its much larger national set of data, consisting of 3,695,943 people. *Id.* at 453.

159. *Id.* at 448.

160. *Id.* at 449 (emphasis added).

per statistical life (“VSL”) is \$8.7 million.¹⁶¹ Dividing the VSL by one half to loosely account for the age and lower health statuses of the patients at issue amounts to \$761,250,000 for the total approximate costs when analyzing the data of this one empirical study.¹⁶² Given the broad use of the HRMT, it is highly likely that the harms that result from AI bias are widespread, numerous, and costly.

4. Extent to Which the Activity Is Not a Matter of Common Usage

Analyzing the “common usage” factor in the AI context is more complicated than it may initially appear. One of the justifications for the existence of the commonality factor is that the law should respect basic public attitudes about the risks that a society is willing to accept:

Basic public attitudes tend to be accepting of familiar and traditional risks, even while apprehensive of risks that are uncommon and novel. The law should be respectful of public attitudes of this sort. When an activity has moved beyond its initial stages and has become common and normal, this tends to allay concerns as to the acceptability of the activity itself.¹⁶³

First, although routine use of AI-enabled tools may generally be well understood, the public is unaware of the more discrete ways that AI is used in various industries, including healthcare.¹⁶⁴ Because society is unaware of the undisclosed uses of AI in certain settings, it cannot broadly consent to its use or the associated risks.

Second, the law permits that when an activity itself is normal, but the unusual risk it creates under particular circumstances is abnormal, the activity may qualify as “abnormally dangerous.”¹⁶⁵ This alternative method for determining commonality opens the door to AI-enabled tools that create unusual risks under particular circumstances being considered “abnormally dangerous.”

AI systems, including HRMTs, are widely used. An activity of common use is one that “is customarily carried on by the great mass of mankind or by many people in the community.”¹⁶⁶ The HRMT is used in hospitals across the nation and is “one of the largest and most typical examples of a class of commercial risk-prediction tools that, by industry estimates, are applied to roughly 200 million people in the United States each

161. INDUS. ECON. INC., CONSUMER PROD. SAFETY COMM’N, FINAL REPORT: VALUING REDUCTIONS IN FATAL RISKS TO CHILDREN 1 (2018), https://www.cpsc.gov/s3fs-public/VSL_Children_Report_FINAL_20180103.pdf [<https://perma.cc/3AAF-5GV5>]. See generally Lisa A. Robinson & James K. Hammitt, *Research Synthesis and the Value per Statistical Life*, 35 RISK ANALYSIS 1086 (2015). Various regulatory agencies determine the VSL, including the Consumer Product Safety Commission, the Environmental Protection Agency, and the Department of Transportation, *inter alia*. Lisa A. Robinson, *Cost-Benefit Analysis and Well-Being Analysis?*, 62 DUKE L.J. 1717, 1724 (2013). VSL estimates have ranged from \$1 million to \$20 million. *Id.*

162. $8,700,000/2 \times 175 = 761,250,000$. See INDUS. ECON. INC., *supra* note 161, at 7.

163. RESTATEMENT (THIRD) OF TORTS: PHYS. & EMOT. HARM § 20 cmt. j (AM. L. INST. 2010).

164. See generally Bernard Marr, *The 10 Best Examples of How AI Is Already Used in Our Everyday Life*, FORBES (Dec. 16, 2019, 12:13 AM), <https://www.forbes.com/sites/bernardmarr/2019/12/16/the-10-best-examples-of-how-ai-is-already-used-in-our-everyday-life/> [<https://perma.cc/82BA-NZ2G>]; Ross & Herman, *supra* note 109.

165. RESTATEMENT (SECOND) OF TORTS § 520 cmt. f (AM. L. INST. 1977).

166. *Id.* § 520 cmt. i.

year.”¹⁶⁷ Increasingly, predictive algorithms have come to shape modern life. Algorithms are used to determine credit scores; they are used to determine rental decisions; they are used to predict worker and student productivity; and they are used to predict the risk of future criminal behaviors to determine pretrial release and sentencing decisions, *inter alia*.¹⁶⁸

Though AI is common, people often do not know if, when, or how frequently they engage with it.¹⁶⁹ This lack of awareness negates ostensible markers indicating social acceptance. A recent Pew Research study showed that although Americans are aware of various ways they may come into contact with AI in their daily lives, only 30% could correctly identify all six examples of AI used in daily life when they were asked about them in a survey.¹⁷⁰ Though 44% of Americans think that they interact with AI less than once per day, most interact with it numerous times per day.¹⁷¹

Most concerning, however, is when AI is used to make vital decisions without the public’s knowledge or consent.¹⁷² For example, AI is used by healthcare insurers “secretly” to determine when they can stop payment for older patients’ Medicare treatment.¹⁷³ Dolores Millam required surgery for a broken leg.¹⁷⁴ After surgery, she was transferred to a nursing home to recover.¹⁷⁵ Her doctor advised her to stay off of her leg for at least six weeks.¹⁷⁶ Unbeknownst to her, Millam’s insurer used an algorithm to predict that she would only need to stay in the nursing home for fifteen days.¹⁷⁷

Soon after the fifteen days, Millam’s covered care was terminated.¹⁷⁸ Millam could not walk or use the facilities without assistance.¹⁷⁹ She had to be moved by mechanical lifts and required around-the-clock care to help her with daily tasks.¹⁸⁰ Millam’s nurse noted that she was “not yet safe to live independently.”¹⁸¹ Millam appealed to Medicare twice.¹⁸² Both appeals were denied.¹⁸³ Millam later filed a federal lawsuit to seek reimbursement for her out-of-pocket costs due to the termination of coverage.¹⁸⁴ The judge found in her favor, saying that the insurance company was not justified in denying

167. Obermeyer et al., *supra* note 92, at 447.

168. Crystal S. Yang & Will Dobbie, *Equal Protection Under Algorithms: A New Statistical and Legal Framework*, 119 MICH. L. REV. 291, 293–94 (2020).

169. Hyken, *supra* note 29; *see also* Obermeyer et al., *supra* note 92, at 447–53.

170. Brian Kennedy, Alec Tyson & Emily Saks, *Public Awareness of Artificial Intelligence in Everyday Activities*, PEW RSCH. CTR. (Feb. 15, 2023), <https://www.pewresearch.org/science/2023/02/15/public-awareness-of-artificial-intelligence-in-everyday-activities/> [https://perma.cc/L33R-HTBT].

171. *Id.*; *see also* Marr, *supra* note 164.

172. *See* Ross & Herman, *supra* note 109.

173. *Id.*

174. *Id.*

175. *Id.*

176. *Id.*

177. *Id.*

178. *Id.*

179. *Id.*

180. *Id.*

181. *Id.*

182. *Id.*

183. *Id.*

184. *Id.*

coverage when the patient was facing “safety risk[s].”¹⁸⁵ The sole reason for Millam’s insurance denial was the insurance company’s reliance on an algorithm to determine when a patient with her health history and injury should be discharged.¹⁸⁶

But for litigation, Millam would never have known that AI was used in the first instance, or that it was the determining factor regarding whether she would continue to receive care or not. Millam’s case is just one example of AI being used as the basis for a crucial decision without an individual’s or the general public’s knowledge or consent. Because society at large does not knowingly consent to the use of AI in many cases, it cannot accept the risks of the widespread harm that it can inflict. In this way, AI differs from other common activities where the public is willing to accept dangerous risks.

In general, for an activity to be abnormally dangerous, it must create a danger of physical harm to others, but it must be an *abnormal* danger.¹⁸⁷ The “unusual risks created by more usual activities under particular circumstances” converts the activity into one that is abnormally dangerous.¹⁸⁸ Indeed, “[t]he essential question is whether the risk created is so unusual, either *because of its magnitude or because of the circumstances surrounding it*, as to justify the imposition of strict liability for the harm that results from it, even though it is carried on with all reasonable care.”¹⁸⁹ This consideration is an important one because AI tools are not abnormal in the frequency of their use—to the contrary, they are widely used across industries.¹⁹⁰ However, the types of *risks* that AI tools create are unusual in the extreme, both because of their magnitude and the circumstances surrounding their use. This is one reason there is widespread concern about AI, with some experts arguing it will lead to human destruction.¹⁹¹ This anxiety is a consequence of the uniqueness of the dangers AI creates.¹⁹² Reflecting on the HRMT example, the dangers associated with its use are unusual almost by definition and result in a selection error that cannot strictly be attributed to lack of reasonable care.

On balance, the “common usage” analysis should be informed by (1) society’s awareness of the AI-enabled tool and its consent to the associated risks and (2) the unusual risks created by the particular circumstances of the activity.

5. Inappropriateness of the Activity to the Place Where It Is Carried On

The inappropriateness of the activity to the place where it is carried on also militates against determining that the HRMT’s use is an abnormally dangerous activity. The HRMT was designed and created to be used in large-scale, high-volume hospital settings,¹⁹³ and that is where it has been deployed. This factor considers that certain activities may be very safe in one environment, but very unsafe in a different

185. *Id.*

186. *Id.*

187. RESTATEMENT (SECOND) OF TORTS § 520 cmt. f (AM. L. INST. 1977).

188. *Id.*

189. *Id.* (emphasis added).

190. See Hyken, *supra* note 29.

191. See HENDRYCKS ET AL., *supra* note 66.

192. See *id.*

193. See Obermeyer et al., *supra* note 92, at 447.

environment.¹⁹⁴ When an activity is conducted with no one within its vicinity, it may not be considered abnormally dangerous.¹⁹⁵ However, when that same activity is conducted in a densely populated area, it may be deemed abnormally dangerous.¹⁹⁶

Envision a scenario where the HRMT *is* used in the wrong environment, causing heightened harm. For example, consider the HRMT's use in an urban environment where all of the hospital's patients are Black. Changing the locality of the application is likely to significantly exacerbate the harms associated with the Tool's recommendations. Under this hypothetical, if the hospital treated only Black patients, the likely outcome would simply be that all of the patients would be selected for the program at lower rates, because the tool was trained using national rather than local datasets.¹⁹⁷

This hypothetical underscores the highly fact-specific nature of each AI activity. Simply examining the HRMT shows that whether it is carried on in an inappropriate locality may be dependent on where the hospital is located. Additionally, a hospital serving only Black patients is not far from reality. Although the statistics are not quite as stark as referenced in the hypothetical, "[t]he 5% of hospitals with the highest volume of [B]lack patients cared for nearly half of all elderly [B]lack patients."¹⁹⁸ Also, "in 2010-2011[,] three-quarters of all Black infants . . . were born in just one-quarter of U.S. hospitals."¹⁹⁹ These statistics show that healthcare for Black patients is highly concentrated in a small percentage of hospitals.²⁰⁰ If the HRMT is used in an urban hospital, it might be an inappropriate location, tipping this factor into the abnormally dangerous realm.

6. The Extent to Which Its Value to the Community Is Outweighed by Its Dangerous Attributes

Traditionally, the more valuable the activity is to the community, the less likely it is that it will be regarded as abnormally dangerous.²⁰¹ The *Restatement (Second)* provides that this factor typically applies when the community in which the activity is engaged is specifically dependent upon or "largely devoted" to the dangerous activity.²⁰² Considering the HRMT, it is likely that a court would conclude that the value of this product to the community outweighs its dangerous character. When it works, it works well. Reducing the length of stay for hospitalized patients and lowering their chances of returning to the hospital with risks of fatal outcomes is a tremendous value to the community. This value, arguably, cannot be achieved by human ability alone. For these reasons, the value to the community may be outweighed by the associated harms.

194. RESTATEMENT (SECOND) OF TORTS § 520 cmt. j (AM. L. INST. 1977).

195. *See id.*

196. *Id.*

197. *See Obermeyer et al., supra* note 92, at 451.

198. Ashish K. Jha, E. John Orav, Zhonghe Li & Arnold M. Epstein, *Concentration and Quality of Hospitals that Care for Elderly Black Patients*, 167 ARCHIVES INTERNAL MED. 1177, 1177 (2007).

199. Gracie Himmelstein, Joniqua N. Ceasar & Kathryn E.W. Himmelstein, *Hospitals that Serve Many Black Patients Have Lower Revenues and Profits: Structural Racism in Hospital Financing*, 38 J. GEN. INTERNAL MED. 586, 586 (2022); *see also* Jha et al., *supra* note 198, at 1177.

200. Himmelstein et al., *supra* note 199, at 586; Jha et al., *supra* note 198, at 1177.

201. RESTATEMENT (SECOND) OF TORTS § 520 cmt. k (AM. L. INST. 1977).

202. *Id.*

C. On Balance, the High-Risk Management Tool May Be Considered Abnormally Dangerous When Analyzed Under the Abnormally Dangerous Activities Test

When considering the abnormally dangerous activities factors as applied to the HRMT example, three out of the six factors arguably weigh on the side of being abnormally dangerous: the existence of a high degree of risk of harm to the person, land, or property of another; the likelihood that the harm that results will be great; and the inability to eliminate the risk by the exercise of reasonable care. This indicates that certain AI tools may be considered abnormally dangerous. This Article does not assert that strict liability is applicable to *all* harms associated with the many potential injuries from AI. Instead, the strict liability framework may be appropriate only when the injury results in harm to persons, land, or property—that is, not economic damages only—and only for those harms that may be categorized as abnormally dangerous activities.

III. THE SIX-FACTOR ABNORMALLY DANGEROUS ACTIVITIES TEST IS USEFUL, BUT SHOULD BE REVISED TO FIT THE SPECIAL CHALLENGES ASSOCIATED WITH ARTIFICIAL INTELLIGENCE

The six-factor abnormally dangerous activities test is a useful tool, but it should be revised to fit AI's special qualities. If each factor is given equal weight, the test is an imperfect fit when applied to unprecedented technology. Well-founded legal paradigms can assist in crafting new legal frameworks to suit problems associated with an unprecedented time. Thus, this Article proposes amending the six-factor test to better fit the harms associated with AI, as provided below²⁰³:

Revised Abnormally Dangerous Activities Six-Factor Test, for AI-related Activities:

1. Existence of a high degree of risk of some harm to the person, land or chattels of others
2. Likelihood that the harm that results from it will be great
3. Inability to eliminate the risk by the exercise of reasonable care
- ~~4. Extent to which the activity is not a matter of common usage~~
- ~~5. Inappropriateness of the activity to the place where it is carried on~~
4. Inability to foresee harms associated with its use, even when reasonable care is exercised
5. Extent to which its value to the community is outweighed by its dangerous attributes

This Section will provide support for a new fourth factor and justification as to why the “common usage” and physical location factors should not be considered.

A. The Inability To Foresee Harms Associated with Its Use, Even When Reasonable Care Is Exercised

This proposed factor, *the inability to foresee harms associated with its use, even when reasonable care is exercised*, reaches the heart of the most unique and perplexing AI problem—AI designers do not fully know how their creations will behave once they

203. All recommended additions are underlined. All recommended deletions are struck.

are let loose in the world.²⁰⁴ Foreseeability for some AI-enabled tools is not knowable. AI designers and researchers themselves admit that they neither know what certain AI tools will do, nor how they will do it.²⁰⁵ The inability to know whether decisions and results are foreseeable or not is a unique attribute of AI, unlike any technology before it.²⁰⁶ This quality of being unknowable provides the basis for a revised fourth factor, because it focuses explicitly on a trait that courts should consider when determining liability for makers of AI tools.

If the AI-enabled activity is one where unavoidable risk is not eliminated with the use of reasonable care (the preceding factor), then it is appropriate to move on to this fourth factor in determining whether the harms are foreseeable even when reasonable care is used. If a court finds that the harm was not foreseeable, then this fourth factor should weigh heavily in considering whether the activity is abnormally dangerous. Moreover, if this one factor is met, it would not be necessary for a court to continue to consider the remaining factors. If an activity could not have been conducted with reasonable care because the putative harms were not knowable at the time of the design, then the actor would not be held to a negligence standard.

Machine learning AI can become defective *after* its original design in ways that cannot be explained or predicted at the outset.²⁰⁷ This can be the case even when the AI design was originally created with the care that one would expect, based on the information available at the time of the tool's creation.²⁰⁸ Consider, for example, the hypothetical use of machine learning AI to assist medical professionals in detecting indications of bone cancer on scans.²⁰⁹ Assume this diagnostic technology uses a deep

204. See Judy Wawira Gichoya et al., *AI Recognition of Patient Race in Medical Imaging: A Modelling Study*, 4 LANCET DIGIT. HEALTH e406, e413 (2022); see also *AI Systems Can Detect Patient Race, Creating New Opportunities To Perpetuate Health Disparities*, EMORY UNIV. (May 27, 2022) [hereinafter *AI Systems Can Detect Patient Race*], https://news.emory.edu/stories/2022/05/hs_ai_systems_detect_patient_race_27-05-2022/story.html [<https://perma.cc/YT54-PUHZ>].

205. See Gichoya et al., *supra* note 204, at e412 (“Although an aggregation . . . features could be partially responsible for the ability of AI models to detect racial identity in medical images, we could not identify any specific image-based covariates that could explain the high recognition performance presented here.”); *AI Systems Can Detect Patient Race*, *supra* note 204 (“We don’t know how the machines are detecting race so we can’t develop an easy solution. . . . Just as with human behavior, there’s not a simple solution to fixing bias in machine learning. The worst thing you can do is try to simplify a complex problem.” (quoting Dr. Judy Gichoya)).

206. See Anat Lior, *The AI Accident Network: Artificial Intelligence Liability Meets Network Theory*, 95 TUL. L. REV. 1103, 1110–11 (2021).

207. See generally Gichoya et al., *supra* note 204, at e412; *AI Systems Can Detect Patient Race*, *supra* note 204.

208. See generally Mark A. Geistfeld, *A Roadmap for Autonomous Vehicles: State Tort Liability, Automobile Insurance, and Federal Safety Regulation*, 105 CALIF. L. REV. 1611, 1645 (2017) (“For example, Google (now Waymo) incorporates machine learning into its self-driving cars and ‘has driven almost two million kilometers on public roads with test drivers and has assembled an enormous fund of traffic situations from which its vehicles can learn.’ Rather than relying on a fixed set of behavioral rules (which characterize symbolic artificial intelligence), the operating system ‘learns’ by adapting or changing the program to incorporate newly acquired information about the best way to execute the dynamic driving task.” (footnotes omitted) (quoting ALEXANDER HARS, INVENTIVIO GMBH, TOP MISCONCEPTIONS OF AUTONOMOUS CARS AND SELF-DRIVING VEHICLES 4 (2017), <https://inventivio.com/innovationbriefs/2016-09/Top-misconceptions-of-self-driving-cars.pdf> [<https://perma.cc/MXV3-37HC>])).

209. See generally Gichoya et al., *supra* note 204.

learning neural network to detect cancer in patients.²¹⁰ This hypothetical diagnostic tool was trained solely on hundreds of thousands of bone imaging datasets from racially diverse pools of data, using no other medical information or factors to make its diagnoses.²¹¹ The diagnostic tool is able to detect cancer faster and more accurately than human physicians who analyze the same images.²¹²

However, the diagnostic tool developed an ability to predict with high accuracy the patient's self-reported race *and makes racially disparate determinations based on this fact*. The diagnostic tool erroneously misdiagnoses Black patients at higher rates than other races. The diagnostic tool's designers did not instruct it—or provide it any information—about race and do not know how the tool even determined the patients' races.²¹³ The diagnostic tool's manufacturers were just as baffled as the patients who were misdiagnosed to discover that the device could determine race without being trained on or given any information about patients' racial backgrounds and, moreover, that it made its diagnoses by factoring in these racial determinations.²¹⁴ This outcome is surprising to the designers of the medical scanner, even though they used reasonable care in their design.²¹⁵

This hypothetical is not far from AI's current abilities. In fact, a recent empirical study found that AI was able to determine patients' races based only on its review of the training sets of x-ray images, mammograms, and radiological images of patients, even though it was not instructed to determine race.²¹⁶ The finding is perplexing because even the researchers, including experts in radiology and computer science, do not know how the machines are detecting race.²¹⁷ For example, the researchers tried to determine whether the AI tool was considering “surrogate covariates” such as whether the machine detected the fact that Black people in general have higher bone density than White people do.²¹⁸ However, when correcting for this factor in the research by removing the availability of such information, the machine was able to determine race at a highly accurate rate and, in fact, better than other statistical analyses that were specifically designed to predict race.²¹⁹

The finding is shocking given that the AI tool was not trained on any information that explicitly or implicitly included race. The tool is a black box and made predictions

210. See generally *id.*

211. See generally *id.*; W. Keith Robinson, *Enabling Artificial Intelligence*, 60 HOUS. L. REV. 331, 340 (2022) (“The rise of big data has allowed for mathematical modeling to reach new heights. Big data is the explosion in the quantity of potentially useful data. The advancements made in computing power and storage have made it possible for this field of endeavor to exist.” (quoting W. Keith Robinson, *Artificial Intelligence and Access to the Patent System*, 21 NEV. L.J. 729, 752 (2021))).

212. See generally Gichoya et al., *supra* note 204.

213. *Id.* at e410–13.

214. *Id.*; see also *supra* Part II.D.

215. See generally Gichoya et al., *supra* note 204.

216. See *id.*

217. See *id.* at e406, e412.

218. *Id.* at e406–10.

219. See *id.* at e410–11.

that were wholly unexpected, so much so that even when a team of researchers study the problem, they cannot begin to explain how the machine reached its conclusions.²²⁰

The above scenario raises the question as to whether the hospital administrators or the physicians themselves should be held responsible for the consequences of this technology when it is introduced. However, hindsight is 20/20. Given that the behavior of machine learning AI can be unpredictable, many of the harms associated with its uses are simply not contemporaneously known. Consider the example of the AI tool that unexpectedly predicts race and the HRMT: the problems were only obvious *after* researchers studied them. This is particularly problematic, because it is not feasible for all AI-enabled tools to be empirically investigated as these were. This is due in large part to proprietary protections that attach to most algorithms, including their training data, objectives, and prediction methodologies.²²¹ Yet even when these products are studied, they are still not well understood.²²² Without information about how AI-enabled tools are flawed, physicians and administrators may be blind to any potential problems. Thus, for most AI technologies used, most physicians, developers, coders, and hospital administrators cannot reasonably anticipate harms not yet discovered.

These considerations lead to the conclusion that the proposed factor, the inability to foresee harms associated with its use, even when reasonable care is exercised, is an important element for courts to consider when determining whether an AI activity is abnormally dangerous. The proposed factor addresses an integral feature of AI—that harms may be unknowable and, thus, unforeseeable.

Whether the AI tool designer could have reasonably foreseen the harms that occurred would be a very fact-specific inquiry. For example, in the case of the AI tool that determines race when not instructed to do so, a court might find that a designer could not reasonably foresee that the tool would identify, *sua sponte*, a patient's race and factor that into its diagnosis.

However, the issue with the HRMT is somewhat different. That Tool used patient spending data to select patients for more intensive treatment and resulted in poorer outcomes for Black patients.²²³ In that case, a court might find that harm to Black patients was reasonably foreseeable because the Tool's consideration of healthcare *costs* as a proxy for healthcare *needs* is fraught with predictable problems.²²⁴

If a court finds that an AI-enabled tool is not conducive to reasonably foreseen harms, then it would be more likely to consider the product an abnormally dangerous one. If so, arguably, courts would not need to weigh the remaining factors of the "abnormally dangerous" test. In contrast, if a court finds that a designer can reasonably foresee the harms associated with the use of the AI tool, then a court would be less likely to consider the activity one that is abnormally dangerous. If the activity is one that is reasonably foreseeable, it may be better suited to a negligence theory of tort law.

220. See *id.* at e412–13.

221. Obermeyer et al., *supra* note 92, at 447; see also Ebrahim, *supra* note 52, at 216–17.

222. See generally Gichoya et al., *supra* note 204.

223. See *supra* Part III.A for a discussion of the tool and resulting outcomes.

224. See Kristin N. Johnson, *Automating the Risk of Bias*, 87 GEO. WASH. L. REV. 1214, 1220 (2019).

1. At the Time of Its Emergence, Nuclear Energy Was Also Deemed Unpredictable, and Therefore Abnormally Dangerous

Nuclear energy was also viewed as unpredictable when it was being developed. The first nuclear reactor in the world was constructed in 1942 on an old squash court at the University of Chicago.²²⁵ In 1946, Congress enacted the Atomic Energy Act, stating in its first Declaration of Policy:

The effect of the use of atomic energy for civilian purposes upon the social, economic, and political structures of today cannot now be determined. It is a field in which unknown factors are involved. Therefore, any legislation will necessarily be subject to revision from time to time. It is reasonable to anticipate, however, that tapping this new source of energy will cause profound changes in our present way of life. Accordingly, it is hereby declared to be the policy of the people of the United States that, subject at all times to the paramount objective of assuring the common defense and security, the development and utilization of atomic energy shall, so far as practicable, be directed toward improving the public welfare, increasing the standard of living, strengthening free competition in private enterprise, and promoting world peace.²²⁶

Justice Byron White of the Supreme Court said, “To facilitate [atomic energy] development the Federal Government . . . erected a complex scheme to promote the civilian development of nuclear energy, while seeking to safeguard the public and the environment from the unpredictable risks of a new technology.”²²⁷ The general view of nuclear power at the time of its emergence was that it was “too complicated, or secret, or mysterious” for lay people to comprehend.²²⁸ This lack of understanding of the then-new nuclear technology in the mid-twentieth century is strikingly similar to the way the public currently views AI. The unpredictability of the harms associated with atomic energy is the predominant reason the production of atomic energy was determined to be an abnormally dangerous activity.²²⁹

While the emergence of atomic energy and AI share common features, AI actually has no *pure* analogue. Although, like all new technology, atomic energy was unknown when it emerged, it did not remain that way. The concept of atomic energy has become “clean, [and] renewable,” as well as “familiar, [and] well-established” and may have greater acceptance than wind energy given its relatively safe history.²³⁰ As its properties have become better understood, nuclear energy is a predictable power source that has resulted in surprisingly few large nuclear occurrences.²³¹

225. OFF. OF NUCLEAR ENERGY, SCI. & TECH., U.S. DEP’T OF ENERGY, THE HISTORY OF NUCLEAR ENERGY 6 (1995), https://www.energy.gov/sites/prod/files/The%20History%20of%20Nuclear%20Energy_0.pdf [<https://perma.cc/9RHF-KW39>].

226. Atomic Energy Act of 1946, Pub. L. No. 79-585, § 1(a), 60 Stat. 755, 755–56 (current version at 42 U.S.C. § 2011).

227. *Pac. Gas & Elec. Co. v. State Energy Res. Conservation & Dev. Comm’n*, 461 U.S. 190, 194 (1983).

228. OFFICE OF THE SEC. OF DEF., ARMED FORCES TALK 276: WHAT TO DO IN ATOMIC ATTACK 6 (1949).

229. See RESTATEMENT (SECOND) OF TORTS § 520 cmt. h (AM. L. INST. 1977).

230. Ashley Hardy & Dontan Hart, Comment, *Policy Meltdown: How Climate Change Is Driving Excessive Nuclear Energy Investment*, 24 BUFF. ENV’T L.J. 137, 181, 183 (2018).

231. See Katherine J. Middleton, Comment, *Danger in 12,008 A.D.: The Validity of the EPA’s Proposed Radiation Protection Standards for the Yucca Mountain Nuclear Repository*, 58 CASE W. RES. L. REV. 933,

In contrast, it is unlikely that AI will become more understood or predictable over time.²³² As advances in machine learning accelerate by servers' ever-increasing power and storage capabilities, algorithms become more sophisticated, and as more data becomes available, AI will become even more difficult for the public to understand.²³³ For example, a new type of AI technology, neuromorphic processing, is more complex than other AI types previously available.²³⁴ Its neural circuits are designed to copy brain processing by working simultaneously (like a real human brain) instead of consecutively.²³⁵ This "neuromorphic cortical model" is likely to result in even more intelligence and higher performance than what is currently available.²³⁶ This means that the inability to reasonably foresee certain outcomes will increase, not decrease. Thus, it is incumbent on courts and legislatures to recognize this fact and create rules to, at the minimum, weigh this as a factor among many when determining liability.

B. Eliminating the "Extent to Which the Activity Is Not a Matter of Common Usage" Requirement

The "extent to which the activity is not a matter of common usage"²³⁷ should be removed from the abnormally dangerous activities test as applied to AI activities. The factor is outdated and not suited to the new challenges associated with AI technology. The *Restatement* is a retrospective summary based on case law developed to that point. Common usage is not relevant to the inquiry because the public often does not know when or how often AI is in use.²³⁸

Moreover, there are instances where courts have found an activity to be abnormally dangerous, even when it is common.²³⁹ For example, the *Restatement (Second) of Torts* provides that if there is physical harm to a person, land, or property on the ground, which was caused by an aircraft taking flight, descending, or falling from the sky, the operator and/or manufacturer is strictly liable, even when the operator and/or manufacturer used

936 (2008). It is worth noting that there are serious risks associated with nuclear waste, which is a problem without readily available or realistic solutions. *Id.* at 936–37. Currently, the United States stores approximately 50,000 metric tons of nuclear waste in 131 temporary cooling pools across the nation. *Id.* at 937. This storage is risky and could put millions of people at risk in the event of pool warming. *See id.*

232. See Anjanette H. Raymond, *Information and the Regulatory Landscape: A Growing Need To Reconsider Existing Legal Frameworks*, 24 WASH. & LEE J. C.R. & SOC. JUST. 357, 377 (2018) ("Currently, the inner workings of an advanced AI is almost impossible to understand, and no amount of transparency will remedy this predicament."); Matthijs M. Maas, *International Law Does Not Compute: Artificial Intelligence and the Development, Displacement or Destruction of the Global Legal Order*, 20 MELB. J. INT'L L. 29, 53 (2019); Note, *Beyond Intent: Establishing Discriminatory Purpose in Algorithmic Risk Assessment*, 134 HARV. L. REV. 1760, 1760, 1763 (2021); Bonnie Kaplan, *Seeing Through Health Information Technology: The Need for Transparency in Software, Algorithms, Data Privacy, and Regulation*, 7 J.L. & BIOSCIENCES, Jan.–June 2020, at 1, 2.

233. Dustin J. Corbett, *A Premier Paradigm Shift: The Impact of Artificial Intelligence on U.S. Intellectual Property Laws*, 17 LIBERTY U. L. REV. 321, 364 (2023); Van der Made, *supra* note 56.

234. Van der Made, *supra* note 56.

235. *Id.*

236. *Id.*

237. RESTATEMENT (SECOND) OF TORTS § 520(d) (AM. L. INST. 1977).

238. *See supra* Part III.B.4.

239. *See* DAN B. DOBBS, PAUL T. HAYDEN & ELLEN M. BUBICK, DOBBS' LAW OF TORTS § 443 (2d ed. 2023).

reasonable care to prevent it.²⁴⁰ The *Restatement (Second)*'s explanation of the rationale is useful:

The position taken is that aviation has not yet reached the stage of development where the risks of accidental physical harm to persons or to land or chattels on the ground is properly to be borne by those who suffer the harm, rather than by the industry itself. The risk of harm to those on the ground is sufficiently obvious if anything goes wrong with the flight; and while the safety record is greatly improved it still cannot be said that the danger of ground damage has been so eliminated or reduced that the ordinary rules of negligence law should be applied. Although there will be relatively few cases in which an airplane falls upon a house, for example, the gravity of the harm resulting when a few tons of flaming gasoline descend upon a dwelling is still a factor to be taken into account. Together with this is the obvious fact that those on the ground have no place to hide from falling aircraft and are helpless to select any locality for their residence or business in which they will not be exposed to the risk, however minimized it may be.²⁴¹

Similar conclusions can be drawn concerning the use of AI. First, given its newness, AI has not reached the point of development where the risks of accidental physical harm are properly borne by those who suffer the harms, rather than by the AI companies. Second, the risk of harm to those end users (e.g., the patients who are subjected to the HRMT that will determine their level of treatment) is sufficiently obvious, and the risk of harm is not eliminated or sufficiently reduced such that ordinary negligence should be applied.²⁴² Third, the gravity of the harm resulting from AI-related incidents can be disastrous, including injury potentially resulting in death.²⁴³ Finally, the harms associated with AI-related incidents are likely to be those in which there is no opportunity to "hide" from the risk of harm.²⁴⁴ AI is widely used, but its wide use is not generally known. One cannot hide from a risk if she is unaware that she is undertaking it.

The uniqueness of AI requires a new set of rules. Given the capabilities of AI-enabled tools, and the likelihood that society will rely on them—even with their associated risks of harm—the extent to which the activity is not a matter of common usage is irrelevant.²⁴⁵

C. *Eliminating the "Inappropriateness of the Activity to the Place Where It Is Carried On" Requirement*

The "inappropriateness of the activity to the place where it is carried on"²⁴⁶ should be eliminated from the abnormally dangerous activities test as applied to AI activities, because, like the common usage factor, it is an outdated and ill-fitting criterion for the

240. RESTATEMENT (SECOND) OF TORTS § 520A (AM. L. INST. 1977).

241. *Id.* § 520A cmt. c.

242. *See supra* Part III.B.

243. *See supra* Part III.B for a discussion of the potential dangers associated with the High-Risk Management Tool when considering the lack of treatment for Black patients admitted to the hospital.

244. *See Hyken, supra* note 29.

245. Indeed, one scholar argues that the uncommon activities requirement was never relevant but was created by the *Restatement (First)* authors in 1938 to avoid hindering dangerous activities that were common and perceived as important to society. Shavell, *supra* note 141, at 39–41.

246. RESTATEMENT (SECOND) OF TORTS § 520 (AM. L. INST. 1977).

advances in technology where the harms themselves are abnormal. Traditionally, “[t]he more appropriate an activity is to its setting, the less likely it is to be considered abnormally dangerous.”²⁴⁷ When considering explosives used at a mine, for example, this factor is well-suited. But it is ill-suited for AI, where physical location has minimal relation to the increased risk of harm.

Moreover, the *Restatement (Second)* provides an alternative path which strikes at the heart of this factor—to consider whether the risk is so unusual, due to its magnitude or the uncommon circumstances around it, that the imposition of strict liability is appropriate.²⁴⁸ As the *Restatement* specifies, “[A]bnormal dangers arise from activities that are in themselves unusual, or from unusual risks created by more usual activities under particular circumstances.”²⁴⁹ Because the magnitude and circumstances related to certain AI-related harms are so uncommon as to be a matter of first impression, the atypicality of the harms are inherent.²⁵⁰ Assuming the ubiquity of AI, strict liability should still apply because of the unusual risks it poses. *Merriam-Webster* defines “abnormal” simply as “deviating from the normal or average.”²⁵¹ The risks of harms associated with certain types of AI may, thus, be deemed abnormal.²⁵² Thus, the inappropriateness of the activity to the place where it is carried out is not a necessary condition for determining abnormality in the context of AI-related risks. Unlike explosives used in a mine, AI is everywhere, though most people are unaware of this fact.

Because the six-factor test for abnormally dangerous activities is imperfect for analyzing the dangers associated with AI, it should be revised to better shape tort law for the emerging issues predominant in society today.

IV. UMBRELLA PROTECTIONS CONSISTENT WITH THE PRICE-ANDERSON ACT WOULD SERVE TO SHIELD CONSUMERS WHILE ENCOURAGING THE AI INDUSTRY TO FLOURISH

The proposed abnormally dangerous activities test may be a useful legal framework to hold companies that design and distribute AI accountable for the risks of harms associated with the use of certain types of AI. However, this approach under a traditional model contemplates protracted litigation against sometimes multibillion-dollar companies with plaintiff consumers receiving relief only in the distant future. A more efficient system of recovery would protect consumers from having to wait to obtain

247. *Schuck v. Beck*, 497 P.3d 395, 419 (Wash. Ct. App. 2021) (citing *Arlington Forest Assocs. v. Exxon Corp.*, 774 F. Supp. 387, 391 (E.D. Va. 1991)).

248. RESTATEMENT (SECOND) OF TORTS § 520 cmt. f (AM. L. INST. 1977); 57A AM. JUR. 2D *Negligence* § 375 (2022) (“As the courts recognize, the essential question is whether the risk created is so unusual, either because of its magnitude or because of the circumstances surrounding it, as to justify the imposition of strict liability for the harm that results from it, even though it is carried on with all reasonable care. In other words, are its dangers and inappropriateness for the locality so great that, despite any usefulness it may have for the community, it should be required as a matter of law to pay for any harm it causes, without the need of a finding of negligence.” (footnotes omitted)).

249. RESTATEMENT (SECOND) OF TORTS § 520 cmt. f (AM. L. INST. 1977).

250. HENDRYCKS ET AL., *supra* note 66.

251. *Abnormal*, MERRIAM-WEBSTER, <https://www.merriam-webster.com/dictionary/abnormal> [https://perma.cc/NYN3-PB3Q] (last visited Apr. 5, 2024).

252. See *supra* Part III.B.

recovery when AI companies are at fault. Moreover, AI victims should not be held financially responsible for shouldering the costs associated with their injuries. AI companies should be held accountable for the harms associated with their technology because even *they* do not fully know what their technology will do and the harms it could create. The burden should be on the AI companies to compensate AI users for the associated harms because by releasing the technology they are, in effect, using consumers as their test subjects. It is for these reasons that a fast-acting policy for damages incurred, combined with a strict liability approach, would be a viable solution for holding AI companies accountable for the harms that they are responsible for.

A. *The Concept of Nuclear Energy Insurance Policies Can Be Applied to the Harms Associated With Artificial Intelligence*

The application of strict liability to the production of atomic energy is the quintessential and most appropriate historical example of a once-new high-risk, high-reward technology.²⁵³ Since the 1950s, the Supreme Court has recognized that investment and growth in the nuclear industry could lead to “potentially vast liability in the event of a nuclear accident of a sizable magnitude.”²⁵⁴ Nuclear energy was certain to lead to “unrectifiable danger” even when reasonable care was used.²⁵⁵ In the 1978 case, *Duke Power Co. v. Carolina Environmental Study Group, Inc.*, the Supreme Court acknowledged that this potentially catastrophic liability could not be adequately addressed by the safety net of insurance.²⁵⁶ In a 1956 congressional hearing, a nuclear industry spokesperson told lawmakers that private operators of nuclear power plants “would be forced to withdraw from the field if their liability w[as] not limited by appropriate legislation.”²⁵⁷ Congress realized that the private insurance available to nuclear operators was “insufficient . . . to cover potential damages from a catastrophic accident. Federal indemnity was therefore considered appropriate to supplement that insurance to assure adequate compensation to the public in the event of a major nuclear accident.”²⁵⁸ In response to these concerns, Congress passed the Price-Anderson Act, in 1957, to “protect the public and to encourage the development of the atomic energy industry, in the interest of the general welfare and of the common defense and security.”²⁵⁹

253. See *Carolina Env’t Study Grp., Inc. v. U.S. Atomic Energy Comm’n*, 431 F. Supp. 203, 223–24 (W.D.N.C. 1977), *rev’d sub nom. Duke Power Co. v. Carolina Env’t Study Grp., Inc.*, 438 U.S. 59 (1978).

254. *Duke Power Co.*, 438 U.S. at 63–64 (providing a thorough history of atomic energy development).

255. *Bennett v. Mallinckrodt, Inc.*, 698 S.W.2d 854, 868 (Mo. Ct. App. 1985) (stating that even in 1985, “[t]he nuclear industry is unique in its inherent and, at present, unrectifiable danger”).

256. 438 U.S. at 64.

257. *Id.* (citing *Government Indemnity for Private Licensees and AEC Contractors Against Reactor Hazards: Hearings Before the Joint Comm. on Atomic Energy*, 84th Cong. 9, 109–10, 272 (1956) [hereinafter *Government Indemnity Hearings* 1956]). The nuclear industry insisted that the risks of large liability as a result of a nuclear event was low. However, the industry threatened to withdraw from the field if Congress did not resolve to find a solution to their exposure and cap liability. *Id.* (citing *Government Indemnity Hearings* 1956, *supra*, at 9, 109–10, 115, 120, 136–37, 148, 181, 195, 240).

258. S. REP. NO. 100-70, at 14 (1988), *reprinted in* 1988 U.S.C.C.A.N. 1424, 1426.

259. 42 U.S.C. § 2012. This was an amendment to the Atomic Energy Act of 1954. 1 KAREN A. GOTTLIEB, *TOXIC TORTS PRAC. GUIDE* § 9:2 (2023). The Atomic Energy Act of 1946 preceded the Price-Anderson Act. Atomic Energy Act of 1946, Pub L. No. 79-585, ch. 724, 60 Stat. 755. The Atomic Energy

The Price-Anderson Act requires that nuclear operators assume liability for court awards to address injuries that result from “extraordinary nuclear occurrences.”²⁶⁰ The Price-Anderson Act permits the Nuclear Regulatory Commission (“NRC”), the body that issues licenses for nuclear reactor operators, to limit liability for nuclear plant operators and set a limit on private insurance for operators.²⁶¹ The Price-Anderson Act also provides a guaranteed pool of compensation to those who are injured in a qualifying nuclear event, in a method that can be paid out more efficiently than would be permitted under traditional tort models.²⁶²

The Price-Anderson Act established two types of insurance for nuclear operators to cover these costs.²⁶³ First, each nuclear operator is required to be covered under the “maximum liability insurance commercially available.”²⁶⁴ Second, any costs from a nuclear event that exceed the amount of publicly available insurance coverage are split equally among all qualified nuclear operators, with a cap.²⁶⁵ For any injuries exceeding

Act was amended in 1954 to provide licensing of construction, ownership, and operation of commercial nuclear reactors to promote the production of energy. Atomic Energy Act of 1954, Pub. L. No. 83-703, ch. 1073, 68 Stat. 919 (codified as amended at 42 U.S.C. §§ 2011–2281 and in scattered additional sections). The Atomic Energy Act can be traced to earlier agreements negotiated by the Manhattan Engineering District of the U.S. Department of War, which organized the private construction and operation of government nuclear facilities before and during the Second World War. OMER F. BROWN II LAW OFFICE, DEP’T OF ENERGY, LEGISLATIVE HISTORY OF GOVERNMENT CONTRACTOR INDEMNIFICATION UNDER THE PRICE-ANDERSON ACT 10–11 (2021), <https://www.energy.gov/sites/default/files/2021-12/7a.%20Brown%20Attachment%20A.pdf> [https://perma.cc/76T9-DJYE]. The government recognized that private contractors needed indemnification against the abnormally dangerous associated hazards of producing nuclear energy. *Id.* Due to the known and unknown risks, commercial insurance that would typically cover industrial harms was not available. *Id.* As a result, the government indemnified private contractors for the harms associated with damages related to nuclear properties. *Id.* This model served as a precursor to the Atomic Energy Act and the Price-Anderson Act. Taylor Meehan, Note, *Lessons from the Price-Anderson Nuclear Industry Indemnity Act for Future Clean Energy Compensatory Models*, 18 CONN. INS. L.J. 339, 343 (2012).

260. Meehan, *supra* note 259, at 347; *see also* 42 U.S.C. § 2210(n). An “extraordinary nuclear occurrence” is defined as:

any event causing a discharge or dispersal of source, special nuclear, or byproduct material from its intended place of confinement in amounts offsite, or causing radiation levels offsite, which the Nuclear Regulatory Commission or the Secretary of Energy, as appropriate, determines to be substantial, and which the Nuclear Regulatory Commission or the Secretary of Energy, as appropriate, determines has resulted or will probably result in substantial damages to persons offsite or property offsite.

42 U.S.C. § 2014(j).

261. 42 U.S.C. § 2210; *see also* Daniel W. Meek, Note, *Nuclear Power and the Price-Anderson Act: Promotion over Public Protection*, 30 STAN. L. REV. 393, 398 (1978); MARK HOLT, CONG. RSCH. SERV., IF10821, PRICE-ANDERSON ACT: NUCLEAR POWER INDUSTRY LIABILITY LIMITS AND COMPENSATION TO THE PUBLIC AFTER RADIOACTIVE RELEASES (2024), <https://crsreports.congress.gov/product/pdf/IF/IF10821> [https://perma.cc/8X6T-CQH2].

262. *See* Meek, *supra* note 261, at 410.

263. 42 U.S.C. §§ 2201–2297h-13; HOLT, *supra* note 261, at 1.

264. § 2210(b); HOLT, *supra* note 261, at 1. There is one company that provides nuclear insurance: American Nuclear Insurers. OFF. OF PUB. AFFS., U.S. NUCLEAR REGUL. COMM’N, NUCLEAR INSURANCE AND DISASTER RELIEF 1 (2022), <https://www.nrc.gov/docs/ML0327/ML032730606.pdf> [https://perma.cc/W8Y7-J5RZ]. The typical annual premium for a one-unit reactor is approximately \$1 million. *Id.*

265. OFF. OF PUB. AFFS., *supra* note 264. Covered nuclear operators must have at least 100 megawatt (100,000 kilowatts) rated capacity, which includes all commercial reactors operating in the United States. § 2210(b)(1)(A); HOLT, *supra* note 261, at 1.

the coverage, Congress would indemnify nuclear operator licensees up to \$500 million.²⁶⁶ This total coverage for each incident amounts to over \$13 billion.²⁶⁷ Notably, to date, no nuclear injuries have penetrated the first layer of insurance coverage since the Price-Anderson Act was established over sixty years ago.²⁶⁸

To ensure a common standard for liability in the event of a nuclear incident, Congress created an amendment to the Price-Anderson Act which required nuclear operators to waive defenses.²⁶⁹ This was thought to be a better approach than enacting a federal statute requiring strict liability, because it would require less federal interference with state law.²⁷⁰ State law often required parties to bring claims related to nuclear incidents under a negligence theory of liability.²⁷¹ A 1966 amendment to the Price-Anderson Act assured that all claims arising out of an “extraordinary nuclear occurrence” would be brought in federal courts instead of state courts, while still permitting state law claims.²⁷² This amendment permitted more predictability and uniformity, given the variation in state laws and courts.²⁷³ The Price-Anderson Act served as a carrot to encourage nuclear operator entrepreneurs and as a protection to the consuming public.²⁷⁴ Indeed, the “paramount policy” of strict liability is “the protection of otherwise defenseless victims of manufacturing defects and the spreading throughout society of the cost of compensating them.”²⁷⁵

The Price Anderson Act has been amended numerous times since its inception, including in 1988.²⁷⁶ A Senate Report from the 1988 amendment summarizes Congress’s intent:

The Price-Anderson system is a comprehensive, compensation-oriented system of liability insurance for Department of Energy contractors and Nuclear Regulatory Commission licensees operating nuclear facilities. *Under the Price-Anderson system, there is a ready source of funds available to compensate the public after an accident, and the channeling of liability to a single entity and waiver of defenses insures [sic] that protracted litigation will be avoided.* That is, the Price-Anderson Act provides a type of “no fault” insurance, by which all liability after an accident is assumed to rest with the facility operator, even though other parties (such as subcontractors or

266. § 2210(c)–(d).

267. HOLT, *supra* note 261, at 1.

268. The closest was the disastrous Three Mile Island 2 Reactor incident, which occurred in 1979. *Id.* at 2. This resulted in liability in the amount of \$71 million. *Id.*; OFF. OF PUB. AFFS., *supra* note 264, at 2.

269. *Duke Power Co. v. Carolina Env’t Study Grp., Inc.*, 438 U.S. 59, 63 (1978); *Inst. of Nuclear Power Operations v. Cobb Cnty. Bd. of Tax Assessors*, 510 S.E.2d 844, 847 (Ga. Ct. App. 1999); S. REP. NO. 100-70, at 15 (1988), *reprinted in* 1988 U.S.C.C.A.N. 1424, 1427–28.

270. *Duke Power Co.*, 438 U.S. at 65–66.

271. 156 AM. JUR. *Trials* § 3 (2018); *McMunn v. Babcock & Wilcox Power Generation Grp., Inc.*, 869 F.3d 246, 281 (3d Cir. 2017) (McKee, J., concurring).

272. 156 AM. JUR. *Trials* § 3 (2018).

273. *Id.*

274. GOTTLIEB, *supra* note 259.

275. *Price v. Shell Oil Co.*, 466 P.2d 722, 725–26 (Cal. 1970).

276. CTR. FOR NUCLEAR SCI. & TECH. INFO., AM. NUCLEAR SOC’Y, THE PRICE ANDERSON ACT: BACKGROUND INFORMATION 3 (2005), <https://wx1.ans.org/pi/ps/docs/ps54-bi.pdf> [<https://perma.cc/Y35J-2PM3>].

suppliers) might be liable under conventional tort principles. . . . If damages exceed the limit established by law, the Price-Anderson Act would require Congress to review the situation and determine what action should be taken to make additional funds available to compensate the public.²⁷⁷

The policy behind the Price-Anderson Act's insurance structure, as reflected in the Senate Report comment, has remained consistent in the face of numerous amendments.²⁷⁸ The Price-Anderson Act's creators understood that those affected by such extraordinary nuclear incidents would need immediate assistance.²⁷⁹ As such, emergency funds collected from the insurance pools could be used for immediate care following an event.²⁸⁰

1. A Two-Tiered Insurance Approach Protects the Consuming Public While Limiting Liability

Under the Price-Anderson Act, plaintiffs may sue the defendant nuclear operators for harms resulting from a nuclear incident.²⁸¹ The purpose of the Act is to create a "no-fault insurance scheme which spreads the cost of an accident uniformly over the entire nuclear industry."²⁸² The first layer of insurance covers up to \$500 million per licensed reactor.²⁸³ The second layer of insurance, for harms that exceed this amount, comes from pooled contributions from reactor operators.²⁸⁴ Each operator is liable for \$131 million per licensed reactor.²⁸⁵ This second-layer onetime payment is called a "retrospective premium."²⁸⁶ This retrospective premium is palatable to industry leaders and would be delivered in prompt up-front payment in the event of a nuclear incident because it is capped.²⁸⁷ Thus, as the number of nuclear operators increases, so too does the available coverage for each nuclear incident.²⁸⁸ In addition to these insurance tiers, Congress, under the Atomic Energy Act, enabled federal indemnification of damages up to another \$500 million for each nuclear incident.²⁸⁹

Congress acknowledged in 1956 that it did not know what the actual liability would be, given the lack of understanding of the atomic energy technology.²⁹⁰ As a result,

277. S. REP. NO. 100-70, at 14 (1988), *reprinted in* 1988 U.S.C.C.A.N. 1424, 1426–27 (emphasis added).

278. *See id.* at 14–15.

279. *See id.*

280. *See id.* at 15.

281. O'Connell, *supra* note 9, at 336.

282. *Price-Anderson Act Amendments of 1987: Hearing on S. 44 and S. 843 Before the Subcomm. on Nuclear Reg. of the Comm. on Env't & Pub. Works*, 100th Cong. 8 (1987) [hereinafter *Price-Anderson Act Amendments Hearing 1987*] (statement of James K. Asselstine, Comm'r, NRC).

283. 10 C.F.R. § 140.11(a)(4) (2024).

284. *Id.*

285. OFF. OF PUB. AFFS., *supra* note 264.

286. *Price-Anderson Act Amendments Hearing 1987*, *supra* note 282, at 1 (statement of Senator John B. Breaux).

287. *Id.* at 10–11 (statement of James K. Asselstine, Comm'r, NRC).

288. Dan M. Berkovitz, *Price-Anderson Act: Model Compensation Legislation? The Sixty-Three Million Dollar Question*, 13 HARV. ENV'T. L. REV. 1, 15 (1989).

289. O'Connell, *supra* note 9, at 339; 42 U.S.C. § 2210(c)–(d).

290. *See Government Indemnity Hearings 1956*, *supra* note 257, at 2 (letter from Sen. Clinton P. Anderson, Chairman).

Congress set the indemnity limit arbitrarily, with an eye towards not “frighten[ing] the country or the Congress to death.”²⁹¹ Although Congress concluded that the \$500 million indemnity umbrella would be sufficient to protect most of the damages that might result from a nuclear incident, its primary focus was on the speed at which plaintiffs could receive compensation.²⁹²

Congress did not intend for plaintiffs to be hamstrung by procedural or other limitations that would slow down the collection of damages to compensate for injuries resulting from nuclear incidents.²⁹³ Therefore, Congress instituted a waiver of defense provision, which would effectuate a strict liability standard for all intents and purposes.²⁹⁴ This means that plaintiffs do not need to prove that the nuclear operators were negligent, instead, they must only prove that the nuclear incident caused their injuries and show the amount of their damages.²⁹⁵

The most significant nuclear incident was the Three Mile Island accident that occurred in 1979 in Pennsylvania.²⁹⁶ Coverage through the Price-Anderson Act was available to injured families and pregnant women who lived close to the Three Mile Island plant.²⁹⁷ In accordance with Congress’s intent that the Act’s remedies be an avenue for speedy recovery of damages, the first-tier insurers advanced funds to the families that needed to pay cost of living expenses related to evacuations.²⁹⁸ In 2003, the insurance available under the Act was also used to settle a class action lawsuit related to the initial incident.²⁹⁹ Additionally, part of the settlement funds were used to create the Three Mile Island Public Health Fund, to study the radiation effects from the incident.³⁰⁰ In total, the available insurance pools covered “\$71 million in claims and litigation costs,” reimbursing more than six hundred parties.³⁰¹

291. *Id.* at 123 (statement of Sen. Clinton P. Anderson, Chairman).

292. *See id.* at 177, 179.

293. *Id.*

294. 42 U.S.C. § 2210(n); 156 AM. JUR. *Trials* § 3 (2018).

295. *See* 156 AM. JUR. *Trials* § 3 (2018); 42 U.S.C. § 2210(n).

296. H. ARCENEUX ET AL., U.S. NUCLEAR REGUL. COMM’N, NUREG/CR-7293, THE PRICE-ANDERSON ACT: 2021 REPORT TO CONGRESS, PUBLIC LIABILITY INSURANCE AND INDEMNITY REQUIREMENTS FOR AN EVOLVING COMMERCIAL NUCLEAR ENERGY OFFICE, 3-1 (2021), <https://www.nrc.gov/docs/ML2133/ML21335A064.pdf> [<https://perma.cc/KUW6-VPCP>] (“The largest nuclear accident that has occurred in the United States took place at a PWR in 1979 at the TMI-2 reactor, near Middletown, Pennsylvania. The TMI-2 accident did not result in any detectable health effects on plant workers or the public. A combination of equipment malfunctions, design-related problems and worker errors led to the partial meltdown of TMI-2 and very small offsite releases of radioactivity” (citation omitted)).

297. *See* OFF. OF PUB. AFFS., *supra* note 264, at 2.

298. *Id.*

299. *Id.*

300. Cassidy S. Davidson, Note, “Hide and Go Seek” Information Policies at the Nuclear Regulatory Commission: Attaining Improved Public Disclosure Could Avert a Nuclear Catastrophe, 45 SW. L. REV. 377, 387 n.77 (2015).

301. OFF. OF PUB. AFFS., *supra* note 264, at 2. One of the most significant objections to the Price-Anderson Act is that it is insufficient to cover the true costs associated with a natural disaster. *See* Jeffrey C. Dobbins, *Promise, Peril, and Procedure: The Price-Anderson Nuclear Liability Act*, 70 HASTINGS L.J. 331, 338–39 (2019). Recognizing this concern, the Nuclear Regulatory Commission requires each nuclear operator licensee to maintain a \$1.06 billion onsite property insurance that covers cleanup costs associated with cleaning

Many critique the Price-Anderson Act's effectiveness and efficiency in addressing damages resulting from nuclear incidents.³⁰² Arguably, though, some protections are better than no protections. Notably, when Congress initially approved the Price-Anderson Act, it acknowledged that liability coverage for nuclear incidents must be "big enough but not too big to scare anyone."³⁰³ The 1988 legislative history related to the Act reveals that at the time of its initial adoption, "it [was] very clear that this [was] a purely arbitrary figure and that no one [had] any real notion—no one had, and to this day we do have a little better notion, but no firm conviction—as to what [a nuclear incident] may cost."³⁰⁴

Further, unchecked free market liability does not mean that compensation in an unregulated sphere would lead to unlimited recovery. Instead, compensation to those harmed would be limited by the "assets of the particular utility involved," which would be set against hardline limits permitted under bankruptcy proceedings.³⁰⁵ Congress recognized the difficulty in determining the amount needed to compensate victims without permitting the nuclear industry to go bankrupt, in the context of the uncertainty of the new technology.³⁰⁶

B. A Combined Strict Liability and Two-Tiered Insurance Policy Approach Would Address Concerns Regarding AI Regulation

In addition to applying strict liability to certain harms associated with AI, the two-tiered insurance approach would prove useful in protecting consumers and AI companies alike. Applying strict liability to certain harms associated with AI would create a fair system for determining liability, where the plaintiff would not need to prove negligence at trial. This would be similar to the nuclear energy industry's mechanism for determining liability. Additionally, a provision similar to what is seen in the Price-Anderson Act would provide swift compensation to those in most need. This is a useful framework because, as with nuclear incidents, some of the potential harms associated with AI are also widespread and disastrous in nature.³⁰⁷

To provide a general proposal, AI companies that currently exist or that wish to enter the AI industry would need to become licensed vendors. To be licensed, an AI manufacturer or designer would need to "buy in" to the industry to obtain a basic level of insurance to protect itself and others from large-scale litigation. The insurance would be fairly conservative in scale, and therefore the premiums would not be prohibitive to startups with limited capital. AI companies would then need to maintain a pool of funds to be used and applied equally to any liability that exceeded the first tier of insurance for any AI company in the industry—with a cap. After the first two tiers of insurance were

a nuclear site and any effected areas. OFF. OF PUB. AFFS., *supra* note 264, at 2. The only company that offers this insurance is Nuclear Electric Insurance Limited. *Id.*; Davidson, *supra* note 300, at 387 n.77.

302. COMPTROLLER GEN., GEN. ACCT. OFF., EMD-80-89, THREE MILE ISLAND: THE FINANCIAL FALLOUT (1980), <https://www.gao.gov/assets/emd-80-89.pdf> [<https://perma.cc/72D4-8CGK>].

303. *Price-Anderson Act Amendments Hearing 1987*, *supra* note 282, at 3 (statement of Sen. George J. Mitchell).

304. *Id.* at 15 (statement of Sen. Alan K. Simpson).

305. *Id.* at 11 (statement of James K. Asselstine, Comm'r, NRC).

306. *See id.*

307. *See supra* Part III.B.

exhausted, the AI companies would not be liable for costs. The amounts for the proposed insurance and the pool of funds are beyond the scope of this Article.³⁰⁸ However, given that there are tens of thousands of AI companies, the pooled liability coverage would be quite significant based even on a conservative estimate.³⁰⁹ This structure would serve to protect the consuming public from waiting for the end of protracted litigation to receive compensation for damages, and it would protect AI companies from going bankrupt, while also providing accountability for the harms associated with their tools.

CONCLUSION

AI tools are augmenting the way that society researches, computes, plans wars, drives, and practices medicine and law. As useful as AI is, it is also a fear-inducing technology that many do not understand well. The fear surrounding AI is not just limited to laypeople. The palpable anxiety surrounding AI is evident in the reticence of legislatures and courts in addressing the question, “How do we create a legal framework for AI?” This question seems almost overwhelming to those who are unfamiliar with the technology, but society must determine an effective answer to this problem.

One feasible approach is to use the well-established law of products liability, specifically, strict liability to govern certain AI-related problems. The “abnormally dangerous activities” test used in the strict liability analysis is useful; however, it cannot adequately address the potential harms resulting from certain types of AI. Thus, this Article recommends a revised “abnormally dangerous activities” test, designed specifically for AI. Further, plaintiffs suing under this theory of liability should be entitled to a speedy recovery, given the wide-reaching effects of some of the potential harms associated with AI. For this reason, Congress should consider creating a two-tiered insurance approach, modeled after the Price-Anderson Act. This would permit swift recovery of damages and limit the time in which plaintiffs could collect payment for the harms incurred. This approach is also useful because it would limit the liability of American AI companies and thus encourage new and continued growth in the industry. History serves as an excellent tutor, but we must also innovate to address the harms that are observed only in modern technology. It is for these reasons that a combined old and novel approach would be most useful.

308. The author is in the process of writing another article on this topic.

309. See *Artificial Intelligence Startups in United States*, TRACXN (March 5, 2024), https://tracxn.com/d/sectors/artificial-intelligence/_cbMnXfS2GfFo4Vi2dxZyUy7l4O8WyzVYLseb9keW5cl/companies [<https://perma.cc/7A2B-WEJ5>].