

ARTIFICIAL INTELLIGENCE AND AGGREGATE LITIGATION

DANIEL WILF-TOWNSEND*

ABSTRACT

The era of AI litigation has begun, and a fundamental tension has emerged. AI tools operate at unprecedented scale, creating mass harms that favor collective legal responses. Yet these same tools generate highly personalized conduct that resists traditional aggregation. This Article argues that the success or failure of AI regulation will be heavily influenced by how well we navigate the tension between scale and personalization in aggregate litigation.

The Article advances a novel framework showing how AI's capacity to both homogenize and differentiate conduct creates opposing forces for class certification. When AI replaces multiple human decision-makers with a single algorithm, it creates common points of failure that facilitate aggregation. But when AI enables mass personalization, it individuates legal claims in ways that may frustrate class treatment. The relative strength of these effects is neither predetermined nor inevitable but will instead be heavily influenced by how we structure AI liability regimes—for instance, by policymakers' decisions about where in the AI supply chain to attach liability and what kind of causal showings to require to prove a claim.

Aggregate litigation can also enable liability regimes for AI harms that might otherwise escape legal remedy. Some AI injuries are statistically demonstrable only at the group level; others arise from aggregate conduct that may cause no actionable individual harm but substantial collective harm. Aggregation creates opportunities to prove harm and assign remedies at the group level, providing a path to address these scenarios. By examining cases and scenarios across contexts including antitrust, antidiscrimination law, and intellectual property, this Article demonstrates how and why policymakers should attend to procedure, and aggregation in particular, as they write the substantive laws governing AI use.

* Associate Professor of Law, Georgetown University Law Center. I would like to thank Maureen Carroll, Ignacio Cofone, Julie Cohen, A. Feder Cooper, Noah Kazis, Alexandra Lahav, Paul Ohm, Alicia Solow-Niederman, Daniel Schwarcz, Nicole Summers, Rachel Wilf-Townsend, and the participants of the CS+Law Workshop, Roundtable on the Law of AI Safety, and annual Civil Procedure Workshop, for sharing thoughts and feedback on the ideas in this Article and earlier drafts. I would also like to thank the editors of the *Washington University Law Review* for their careful attention and thoughtful improvements.

TABLE OF CONTENTS

INTRODUCTION	778
I. ARTIFICIAL INTELLIGENCE AND THE AUTOMATION OF JUDGMENT-INTENSIVE TASKS	784
A. <i>The growth of AI tools and AI harms</i>	784
B. <i>The growing world of AI legislation, regulation, and litigation</i> .	789
II. THE IMPORTANCE OF PRIVATE AGGREGATE ENFORCEMENT FOR AI REGULATION	791
A. <i>Decentralization</i>	793
B. <i>Scale</i>	796
C. <i>Litigation Costs</i>	797
III. THE AVAILABILITY OF AGGREGATION TO MANAGE AI HARMS	800
A. <i>The influence of AI tools on conduct and aggregation</i>	802
1. <i>The homogenizing effect</i>	802
2. <i>The differentiating effect</i>	806
3. <i>The proof effect</i>	808
B. <i>Responses by liability regimes that facilitate aggregation</i>	812
1. <i>Moving liability upstream</i>	812
2. <i>Minimizing causation requirements</i>	816
IV. USING AGGREGATION TO MANAGE PROBLEMS OF GROUP LIABILITY	818
A. <i>Probabilistic wrongs</i>	819
B. <i>Aggregate wrongs</i>	825
CONCLUSION	828

INTRODUCTION

Although the “era of artificial intelligence”¹ is well under way, the era of AI *litigation* is still taking its first steps.² The last decade has seen an explosion of AI tools with a wide range of applications, from medicine to marketing to computer programming.³ But the rapid adoption of this new technology can also lead to serious harms, whether due to bias, inaccuracy,

1. See, e.g., Michael R. Siebecker, *Making Corporations More Humane Through Artificial Intelligence*, 45 J. CORP. L. 95, 101–13 (2019) (discussing “The Era of Artificial Intelligence”); Sylvia Lu, *Algorithmic Opacity, Private Accountability, and Corporate Social Disclosure in the Age of Artificial Intelligence*, 23 VAND. J. ENT. & TECH. L. 99, 100 (2020) (referring to the present, in 2020, as “the age of artificial intelligence”); Alicia Solow-Niederman, *Administering Artificial Intelligence*, 93 S. CAL. L. REV. 633, 644 (2020) (discussing “governance options in the age of AI”).

2. See, e.g., Alicia Solow-Niederman, *Do Cases Generate Bad AI Law?*, 25 COLUM. SCI. & TECH. L. REV. 261, 272–75 (2024) (discussing early litigation against companies making generative AI products).

3. See *infra* Part I.

misuse, or more.⁴ It is therefore perhaps unsurprising that we have recently seen lawsuits about AI tools emerging in areas as diverse as intellectual property, employment discrimination, health insurance, consumer finance, and privacy.⁵

As society grapples with the best way to manage these new tools, many people focus on substantive law. What rules of liability should apply to these kinds of devices?⁶ Will existing common law and statutory regimes map well to the contours of this new technology and its various uses?⁷ While some lawmakers emphasize applying the laws we have,⁸ others are attempting new approaches—including hundreds of bills proposed in federal and state legislatures in the last year alone.⁹

But whatever liability rules we end up with, they will have to be enforced. And the rules governing that enforcement will have significant influence over whether these AI policies succeed or fail.¹⁰

This Article argues that getting AI policy right will require paying particular attention to one of our legal system's most important procedural tools: the class action. The story of AI is in many ways the story of a technology that allows for *scale*. AI tools facilitate scale by automating tasks—whether pricing decisions,¹¹ medical diagnoses,¹² or customer

4. See *infra* Section I.A (describing AI-related harms).

5. See *infra* Parts III & IV (discussing cases).

6. See, e.g., Bryan H. Choi, *AI Malpractice*, 73 DEPAUL L. REV. 301 (2024); W. Nicholson Price II, Sara Gerke & I. Glenn Cohen, *Liability for Use of Artificial Intelligence in Medicine*, in RESEARCH HANDBOOK ON HEALTH, AI AND THE LAW 150 (Barry Solaiman & I. Glenn Cohen, eds. 2024); Mihailis E. Diamantis, *Employed Algorithms: A Labor Model of Corporate Liability for AI*, 72 DUKE L.J. 797 (2023); Andrew D. Selbst, *Negligence and AI's Human Users*, 100 B.U. L. REV. 1315, 1315 (2020); Frank Pasquale, *Data-Informed Duties in AI Development*, 119 COLUM. L. REV. 1917 (2019); David C. Vladeck, *Machines Without Principals: Liability Rules and Artificial Intelligence*, 89 WASH. L. REV. 117 (2014).

7. See, e.g., Charlotte A. Tschider, *Medical Device Artificial Intelligence: The New Tort Frontier*, 46 BYU L. REV. 1551, 1557 (2021) (describing “the current and insufficient state of FDA regulation” for AI tools in medicine); Deborah Hellman, *Measuring Algorithmic Fairness*, 106 VA. L. REV. 811, 846–64 (2020) (examining how the use of racial classifications within AI tools can improve accuracy and reduce algorithmic unfairness in ways consistent with existing antidiscrimination law).

8. See, e.g., CONSUMER FIN. PROT. BUREAU, DEP'T OF JUSTICE, EQUAL EMP. OPPORTUNITY COMM'N & FED. TRADE COMM'N, JOINT STATEMENT ON ENFORCEMENT EFFORTS AGAINST DISCRIMINATION AND BIAS IN AUTOMATED SYSTEMS (2023), https://www.ftc.gov/system/files/ftc_gov/pdf/EEOC-CRT-FTC-CFPB-AI-Joint-Statement%28final%29.pdf [<https://perma.cc/KBH7-JQ6G>].

9. See Grant Gross, *The Complex Patchwork of US AI Regulation Has Already Arrived*, CIO (Apr. 5, 2024), <https://www.cio.com/article/2081885/the-complex-patchwork-of-us-ai-regulation-has-already-arrived.html> [<https://perma.cc/WE69-46YD>].

10. See *infra* Part II.

11. See, e.g., Huiyi Chen, *Insuring Against Antitrust Actions Targeting Pricing Algorithms*, A.B.A. (Mar. 29, 2024), <https://www.americanbar.org/groups/litigation/resources/newsletters/insurance-coverage/insurance-antitrust-pricing-algorithms/> [<https://perma.cc/XSY4-7GAX>].

12. See, e.g., Lynda Seminara, *Smartphone-Based ROP Screening*, EYENET MAG. (July 1, 2023), <https://www.aao.org/eyenet/article/smartphone-based-rop-screening> [<https://perma.cc/BJ4N-XU3E>].

service conversations¹³—that previously required greater human judgment.¹⁴ But in doing so, AI tools also create the potential for correspondingly scaled harms.¹⁵ The class action is the law’s basic answer to the question of how to resolve large numbers of similar legal claims together—making it especially important as AI tools become more widely used and their potential for widespread harm grows.

This Article provides the first general account of the relationship between aggregate litigation and the enforcement of laws governing artificial intelligence.¹⁶ Focusing on class actions, it starts by making the basic argument that aggregate litigation will have an important role to play in any AI regulatory regime that permits private enforcement.¹⁷ In addition to aggregate litigation’s generally important role in U.S. civil justice, there are AI-specific reasons for this to be the case: The scale, decentralization, and complexity of AI tools all point in favor of aggregate litigation as a method of responding to and deterring AI harms.¹⁸

13. See, e.g., Simon Bamberger, Nicholas Clark, Sukand Ramachandran & Veronica Sokolova, *How Generative AI Is Already Transforming Customer Service*, BOS. CONSULTING GRP. (July 6, 2023), <https://www.bcg.com/publications/2023/how-generative-ai-transforms-customer-service> [https://perma.cc/S2UY-YKVA].

14. To be sure, there is a world of hype surrounding AI, and even many notable advances fall short of the pitches of the most zealous boosters. See, e.g., Gerrit De Vynck, *The AI Hype Bubble Is Deflating. Now Comes the Hard Part.*, WASH. POST (Apr. 25, 2024), <https://www.washingtonpost.com/technology/2024/04/18/ai-bubble-hype-dying-money/> [https://perma.cc/8Y9H-XGFQ]. But real improvements in computer science have led to the significant, widespread adoption of many AI tools, and their uptake is rapidly expanding. See, e.g., NESTOR MASLEJ ET AL., INST. FOR HUM.-CENTERED AI, STAN. UNIV., *THE AI INDEX 2024 ANNUAL REPORT* (2024), https://aiindex.stanford.edu/wp-content/uploads/2024/04/HAI_2024_AI-Index-Report.pdf [https://perma.cc/Q7PD-MD29] [hereinafter *AI INDEX 2024*].

15. See *infra* Section I.A.

16. Within the broad contours of “artificial intelligence” and “civil procedure,” there is a robust literature considering how AI tools might be used by lawyers, courts, or arbitrators as a matter of legal practice. See, e.g., David Horton, *Forced Robot Arbitration*, 109 CORNELL L. REV. 679 (2024); Benjamin Minhao Chen, Alexander Stremitzer & Kevin Tobia, *Having Your Day in Robot Court*, 36 HARV. J.L. & TECH. 127, 128 (2022); David Freeman Engstrom & Jonah B. Gelbach, *Legal Tech, Civil Procedure, and the Future of Adversarialism*, 169 U. PA. L. REV. 1001 (2021); Zoe Niesel, *Machine Learning and the New Civil Procedure*, 73 SMU L. REV. 493 (2020). Articles in that literature sometimes focus on class actions. See Alissa del Riego & Joseph Avery, Essay, *The Class Action Megaphone: Empowering Class Members with an Empirical Voice*, 76 STAN. L. REV. ONLINE 1, 5–8 (2023); Peter N. Salib, *Artificially Intelligent Class Actions*, 100 TEX. L. REV. 519 (2022); Omri Ben-Shahar, Response, *Personalized Class Actions*, 100 TEX. L. REV. ONLINE 162 (2022).

This Article, in contrast, focuses on the role of civil procedure in enforcing the regulation of artificial intelligence. Rather than questions like “how will lawyers and courts use AI tools,” this Article addresses questions about how AI regulations in any sector will be enforced via civil litigation. These questions about the relationship between civil procedure and AI regulation are currently a less-explored area, with some notable articles but none focused on class actions. See Solow-Niederman, *supra* note 2; Mark A. Lemley & Bryan Casey, *Remedies for Robots*, 86 U. CHI. L. REV. 1311 (2019).

17. See *infra* Part II.

18. See *infra* Part II.

But resolving claims at scale is not as simple as just adding more names to a case caption. The Article argues that there is a central tension that emerges in the application of law to mass AI harms. AI tools enable both scale and personalization.¹⁹ And while scale makes it easier to aggregate legal claims, personalization makes it harder.²⁰ Understanding and navigating this tension should be a central goal for policymakers who seek to craft enforceable liability regimes.²¹

Whether aggregation is feasible in a particular case is a function of both the facts and the law: The law determines what a plaintiff needs to prove to establish liability, and whether that can be done in the aggregate will depend on the kinds of facts involved in the case. Will the use of AI tools at scale create the kinds of commonalities that provide the foundation for class litigation? Or will the increasingly personalized processing that can now be automated defeat efforts to aggregate claims? These questions have gone largely unaddressed, with most literature around AI regulation focused on substance rather than civil procedure.²²

Focusing on the “facts” part of the equation first, the Article describes several effects of AI tool use that will influence the availability of aggregation.²³ The first is a “homogenizing effect”: Where many different actors or decisionmakers are replaced by a single automated system, that arrangement is likely to facilitate class actions, by making it more likely that there is a common point of failure or other basis for liability across everyone in an affected class.²⁴ Early examples of litigation in which this homogenizing effect is on display include lawsuits in the worlds of rental pricing, consumer finance, and home insurance.²⁵

Conversely, the Article also identifies a “differentiating” effect that some AI tools will have that will make class certification more difficult.²⁶ Technology that enables more personalized and nuanced automation may lead an actor to replace previously uniform conduct with a range of options mediated by an AI tool, as when a company introduces more differentiated products, or replaces information on a static web page with an interactive chatbot.²⁷ These substitutions will result in more heterogeneous treatment. When that varied treatment is material to a set of claims, it will tend to make

19. See *infra* Part I.

20. See *infra* Part III.

21. See *infra* Part IV.

22. See *supra* notes 6–7.

23. See *infra* Section III.A.1.

24. See *infra* Section III.A.1.

25. See *infra* Section III.A.1.

26. See *infra* Section III.A.2.

27. See *infra* Section III.A.2.

class certification more difficult.²⁸ And although AI tools may also provide new ways for courts to manage some kinds of classes, whether courts will use such tools will depend in part on their willingness to use statistical reasoning during class adjudication, which courts have often been hesitant to do.²⁹

Next, the Article turns to the “law” part of the equation governing whether aggregation will be available.³⁰ It identifies two important moves in actual and proposed AI laws that will influence the feasibility of aggregation.³¹ First, lawmakers can move liability “upstream” to actors earlier in a chain of conduct, which increases the likelihood that the material actions will be common to all of those affected by “downstream” conduct.³² So, for instance, holding the developer of an AI tool liable for discrimination or fraud engaged in by an end user will facilitate aggregate litigation by those affected by that discrimination or fraud. Second, lawmakers can lessen the significance of causation requirements for liability. They can, for instance, create “per se” liability regimes that attach liability to the wrongful use of a legally infringing AI tool in the first place rather than requiring proof that using the AI tool caused a particular type of outcome.³³ These strategies, both of which have been implemented in legislation and regulation in at least some contexts, also have their downsides; but a potential benefit of each is that they enable more effective private aggregate enforcement.

Finally, the Article flips the usual script in which procedure is viewed as just a vehicle for substance. It argues that aggregate enforcement may in fact be a key to unlocking new types of liability regimes for governing AI.³⁴ That is because AI tools can cause harms that are difficult to deal with at the level of the individual case—in ways that go beyond the problems of costs and efficiencies that are the usual stock-in-trade of class action conversations.

In some scenarios, for instance, it may be possible to demonstrate statistically that a particular group was harmed—say, that Black content creators as a group were discriminated against in the application of a content moderation algorithm—without being able to prove that any specific individual in that group was harmed.³⁵ Or in other scenarios, it may be that

28. See *infra* Section III.A.2.

29. See *infra* Section III.A.3.

30. See *infra* Section III.B.

31. See *infra* Section III.B.

32. See *infra* Section III.B.

33. See *infra* Section III.B.

34. See *infra* Part IV.

35. See *infra* Section IV.A (discussing “probabilistic wrongs”).

no individual in a group has a clearly valid claim, but that the group in the aggregate has something that looks like a traditional basis for a claim.³⁶ For instance, in pending intellectual property cases against generative AI companies, it is plausible that there are many individual copyright holders whose property is not directly responsible for any trained AI model's skills or outputs.³⁷ As a result, these individuals would have no valid claim for unjust enrichment against the defendants—even if perfect information were available—because their intellectual property is not causally linked to any of the defendant's commercial revenues.³⁸ But it is also plausible that, in the aggregate, all of the individual property owners' IP at issue is strongly causally linked to much of that revenue. In all of these scenarios, aggregate litigation provides a mechanism for a liability regime to address claims that are only colorable in a group setting.

Throughout, the Article aims to show why getting AI governance “right” will require paying significant attention to procedure, and to the availability of aggregate litigation in particular. Class actions have historically been a critical tool for access to justice, enabling consumers, workers, civil rights plaintiffs, and other marginalized individuals to proceed with their claims when courthouse doors would otherwise likely be closed.³⁹ And the harms associated with AI tools are likely to fall on those same groups, with problems of discrimination, deceptive commercial practices, economic injuries, and privacy harms appearing particularly concerning.⁴⁰ Although both the legal procedures and AI products involved in these issues can appear quite technical, the harms are concrete, and significant equities are at stake.

The Article proceeds as follows. Part I provides a brief description of the rapidly increasing use of AI tools in society and the nascent legal response. Part II then discusses the importance of aggregate litigation as part of any AI regulatory regime that permits private enforcement. Next, Part III identifies important influences on the availability of aggregation for AI-related harms, considering the effects of AI on people's conduct as well as the influence of substantive law on the ease of aggregation. Finally, Part IV diagnoses types of wrongs that are likely in the AI context that will be

36. See *infra* Section IV.B (discussing “aggregate wrongs”).

37. See *infra* Section IV.B.

38. See *infra* Section IV.B.

39. See Myriam Gilles, *Class Warfare: The Disappearance of Low-Income Litigants from the Civil Docket*, 65 EMORY L.J. 1531 (2016) (discussing the historic role of class actions in protecting marginalized groups).

40. See, e.g., Alicia Solow-Niederman, *Information Privacy and the Inference Economy*, 117 NW. U. L. REV. 357 (2022); Talia B. Gillis, *The Input Fallacy*, 106 MINN. L. REV. 1175 (2022); Sandra G. Mayson, *Bias in, Bias Out*, 128 YALE L.J. 2218 (2019).

difficult to prove in individual cases, and considers how aggregate litigation can enable new types of liability structures to overcome these challenges.

I. ARTIFICIAL INTELLIGENCE AND THE AUTOMATION OF JUDGMENT-INTENSIVE TASKS

A. The growth of AI tools and AI harms

Over the last decade or so, the use of AI tools in society has exploded.⁴¹ Significant and ongoing advances in computer science have led to major improvements in the field of artificial intelligence.⁴² These improvements, in turn, have resulted in a steep rise in the research, development, and deployment of AI tools, with applications in science, engineering, healthcare, education, government, and more.⁴³

As a result, corporate use of AI has become pervasive in a very short time frame: In 2017, 20% of businesses in a large, cross-industry nationwide survey reported implementing AI in at least one business unit or function;

41. At the outset, a brief note on word choice—especially the phrases “artificial intelligence,” “machine learning,” “algorithm,” and “model.” See David Lehr & Paul Ohm, *Playing with the Data: What Legal Scholars Should Learn About Machine Learning*, 51 U.C. DAVIS L. REV. 653, 669 (2017) (noting that “it has become all too common” for legal scholarship to confuse terms such as “machine learning,” “artificial intelligence,” and “big data analytics”). The use of the phrase “artificial intelligence” in the legal context is broad. See, e.g., 15 U.S.C. § 9401(3) (providing a broad definition of artificial intelligence). “Machine learning,” meanwhile, refers to a more specific family of approaches for building AI tools in which a computerized model is trained on data to “learn” from that data, discovering correlations in the data that build its capacities for prediction and analysis rather than having those capacities directly programmed in. See, e.g., RAYMOND T. NIMMER, JEFF C. DODD & LORIN BRENNAN, 1 INFORMATION LAW § 1:16 (2025) (contrasting rules-based approaches and machine learning approaches to developing AI systems for providing advice and solving problems).

The explosion in AI tools in recent years has been largely due to the result of advances in machine learning. See Jeffrey Dean, *A Golden Decade of Deep Learning: Computing Systems & Applications*, 151 DAEDALUS 58, 62–66 (2022), <https://www.amacad.org/publication/golden-decade-deep-learning-computing-systems-applications> [<https://perma.cc/KHS3-R62B>]. This Article therefore primarily focuses on machine-learning tools, and may use phrases such as “AI tool” or “machine-learning tool” relatively interchangeably, even though it is possible to build AI tools through other means. And it uses the phrase “algorithm” or “model” in a relatively narrow sense, to mean the result of a machine-learning training process: the tool that can take an input and produce an output based on the data that it has been trained on. See, e.g., *Model*, GENLAW GLOSSARY, <https://blog.genlaw.org/glossary.html#model> [<https://perma.cc/3B7V-VYP4>] (describing a model as “a mathematical tool that takes an input and produces an output”). Although this use is more confined than broader, technical definitions of “algorithm” and “model,” it tracks the usage of these terms in many policy conversations. See, e.g., JOINT STATEMENT ON ENFORCEMENT EFFORTS AGAINST DISCRIMINATION AND BIAS IN AUTOMATED SYSTEMS, *supra* note 8 (using the term “algorithm” to refer to automated decision-making tools trained via machine learning).

42. See Dean, *supra* note 41, at 58–62.

43. See *id.* at 62–66.

by 2024, that number was 72%.⁴⁴ For natural persons, in their roles as consumers, employees, and citizens, encounters with AI tools have correspondingly become common: An AI tool might approve or deny a home loan;⁴⁵ set a price for a retail purchase;⁴⁶ flag an individual for increased scrutiny from law enforcement;⁴⁷ determine who gets interviewed or hired for a job;⁴⁸ or advise customers on financial investments.⁴⁹ AI tools have begun to be integrated into industry, commerce, and daily life, and a broad set of indicators suggests that these trends will intensify in the years ahead.⁵⁰

There are a variety of ways to categorize these tools, but an increasingly important dichotomy is between what are often called “generative” AI tools and “predictive” AI tools.⁵¹ Generative AI refers to machine learning tools that generate text, images, or audio; the most broadly familiar example is probably ChatGPT. These tools allow for the automatic drafting of even relatively nuanced documents as well as more sophisticated “live” interactions between individuals and chatbots. Examples of generative AI tools include a legal assistant software that summarizes trial and deposition transcripts,⁵² or a graphic design tool that allows users to input a text

44. Alex Singla, Alexander Sukharevsky, Lareina Yee, Michael Chui & Bryce Hall, *The State of AI in Early 2024: Gen AI Adoption Spikes and Starts to Generate Value*, MCKINSEY & CO. (May 30, 2024), <https://www.mckinsey.com/capabilities/quantumblack/our-insights/the-state-of-ai-2024> [https://perma.cc/TDJ9-AA6T].

45. See Gillis, *supra* note 40, at 1176.

46. See Oren Bar-Gill, Cass R. Sunstein & Inbal Talgam-Cohen, *Algorithmic Harm in Consumer Markets*, 15 J. LEGAL ANALYSIS 1 (2023).

47. See Michael L. Rich, *Machine Learning, Automated Suspicion Algorithms, and the Fourth Amendment*, 164 U. PA. L. REV. 871 (2016).

48. See Pauline T. Kim, *Data-Driven Discrimination at Work*, 58 WM. & MARY L. REV. 857, 860–63 (2017).

49. Daniel Schwarcz, Tom Baker & Kyle Logue, *Regulating Robo-Advisors in an Age of Generative Artificial Intelligence*, 82 WASH. & LEE L. REV. 775 (2025).

50. See generally AI INDEX 2024, *supra* note 14 (describing trends in the investment in and use of AI systems using a wide range of metrics).

51. The terms here are somewhat in flux. Although there seems to be a consensus around using the term “generative AI” to describe tools like large language models, there is less of a consensus around what to call non-generative AI tools, and how exactly to draw the line. Although “predictive AI” is common, some refer to these tools as “traditional AI,” “analytical AI,” “discriminative AI,” and other terms. See, e.g., Katherine Lee, A. Feder Cooper & James Grimmelmman, *Talkin’ ‘Bout AI Generation: Copyright and the Generative-AI Supply Chain*, 72 J. COPYRIGHT SOC’Y 251, 263 (2025) (using “discriminative”); *When to Use Generative AI or Traditional AI*, GOOGLE CLOUD (Sept. 4, 2025), <https://cloud.google.com/docs/ai-ml/generative-ai/generative-ai-or-traditional-ai> [https://perma.cc/M386-47PB] (“traditional”); THOMPSON REUTERS, A GAME-CHANGING SHIFT: WHAT AI MEANS FOR MODERN LAWYERS TODAY 4 (2023), <https://legalsolutions.thomsonreuters.co.uk/content/dam/ewp-m/documents/legal-uk/en/pdf/reports/a-game-changing-shift-tr3718039.pdf> [https://perma.cc/73BH-CW6E] (“analytical”).

52. See, e.g., Jeffrey I. Ehrlich, *A First-Hand Experience with AI Legal Software*, PLAINTIFF MAG., Nov. 2023, <https://plaintiffmagazine.com/recent-issues/item/a-first-hand-experience-with-ai>

description of an image that they would like and then generates a draft of that image.⁵³

“Predictive AI” refers to the use of machine learning tools that make inferences about informational inputs, typically in ways that result in small, discrete outputs that are used for analysis: a “yes/no” binary decision, a recommendation among a limited set of options, or a number.⁵⁴ These tools have broad applications over a large number of industries, including industries where litigation is frequent such as medical care or consumer finance. Examples of predictive AI tools include a multi-cancer screening test trained on thousands of blood samples to recognize cancer signals,⁵⁵ or a software system for insurers that makes pricing and underwriting recommendations.⁵⁶

Although predictive and generative AI have a wide variety of use cases, many uses of both share the important feature of allowing for the increased automation of tasks that have typically involved human attention and judgment. The use of a predictive AI tool to diagnose diabetic retinopathy, for instance, allows for the detection of this sight-threatening disease using a smartphone-mounted camera and software rather than relying on the

legal-software [https://perma.cc/B2EL-3BHX] (summarizing one user’s experience with the CoCounsel research tool).

53. See, e.g., Jess Weatherbed, *Adobe’s New Firefly Model Makes It Easier to Use Photoshop’s AI Tools*, THE VERGE (Apr. 23, 2024, 4:00 AM), https://www.theverge.com/2024/4/23/24138011/adobe-firefly-3-ai-model-photoshop-tools-reference-image [https://perma.cc/D5VG-9XCP].

54. See, e.g., Lee et al., *supra* note 51, at 260–66. The definitions of predictive and generative AI I offer here attempt to track their distinct social and commercial roles, rather than capturing specific underlying technical distinctions. Many legal and policy questions track these social and commercial distinctions described above—with, e.g., a device that makes medical diagnoses raising distinct legal issues compared to a tool designed to produce usable images in response to text prompts. But the technical distinctions between generative AI and predictive AI are also important and may be relevant for a variety of policy issues. See, e.g., A. Feder Cooper et al., *Machine Unlearning Doesn’t Do What You Think: Lessons for Generative AI Policy and Research* 3–4 (unpublished manuscript) https://arxiv.org/abs/2412.06966 [https://perma.cc/W6ML-8W4C] (discussing the relevance of the distinction for machine “unlearning” tools, and the policy relevance of those tools).

55. See, e.g., Janet Vittone, David Gill, Alex Goldsmith, Eric A. Klein & Jordan J. Karlitz, *A Multi-Cancer Early Detection Blood Test Using Machine Learning Detects Early-Stage Cancers Lacking USPSTF-Recommended Screening*, 8 NPJ PRECISION ONCOLOGY, Apr. 17, 2024, art no. 91.

56. See, e.g., Violet Chung, Pranav Jain & Karthi Purushothaman, *Insurer of the Future: Are Asian Insurers Keeping Up with AI Advances?*, MCKINSEY & CO. (May 3, 2023), https://www.mckinsey.com/industries/financial-services/our-insights/insurer-of-the-future-are-asian-insurers-keeping-up-with-ai-advances [https://perma.cc/2FYU-5D8D]; MICHAEL CHUI ET AL., MCKINSEY GLOB. INST., *NOTES FROM THE AI FRONTIER: INSIGHTS FROM HUNDREDS OF USE CASES* (2018), https://www.mckinsey.com/~/media/mckinsey/featured%20insights/artificial%20intelligence/notes%20from%20the%20ai%20frontier%20applications%20and%20value%20of%20deep%20learning/notes-from-the-ai-frontier-insights-from-hundreds-of-use-cases-discussion-paper.pdf [https://perma.cc/2AAC-SNP3].

expertise of an ophthalmologist.⁵⁷ In the world of software coding, generative AI tools have been found capable of handling routine tasks in ways that save between 20% and 50% of coders' time.⁵⁸ AI tools can make automatic pricing decisions across many different items in a business's inventory.⁵⁹ They can have responsive, real-time "conversations" simultaneously with many different customers to convey information and address needs or preferences.⁶⁰ In some circumstances, these tools may substitute for humans entirely, as when an AI diagnostic is used rather than a doctor for a preliminary screening; at other times, they may complement human effort and only partially replace it, as when a coder uses a generative AI tool to complete a low-skill part of a task to free up time for a more complex part of the task.

When AI tools substitute for human attention and judgment, they may be deployed at scales that would be costly or infeasible to achieve with people, or they may simply be used as low-cost replacements for tasks previously done by people. Where a smartphone can be used to diagnose medical conditions, it is possible to create mass screening programs that might otherwise be cost-prohibitive in lower-income nations.⁶¹ Where generative AI allows for chatbots to handle more sophisticated customer service inquiries that usually require human employees to manage, a company may simply hire fewer human employees.⁶² In either of these scenarios, a single AI tool may do work that would previously have taken dozens of individual humans, hundreds, or more.⁶³

With this kind of wide-scale use comes the potential for correspondingly high-volume harms. Diagnostic algorithms can be negligently designed or

57. Ramachandran Rajalakshmi, Radhakrishnan Subashini, Ranjit Mohan Anjana & Viswanathan Mohan, *Automated Diabetic Retinopathy Detection in Smartphone-Based Fundus Photography Using Artificial Intelligence*, 32 EYE 1138 (2018), <https://www.nature.com/articles/s41433-018-0064-9> [<https://perma.cc/J4P2-B8PR>] (finding that an automated detection system was highly sensitive at detecting diabetic retinopathy).

58. Begum Karaci Deniz, Chandra Gnanasambandam, Martin Harrysson, Alharith Hussin & Shivam Srivastava, *Unleashing Developer Productivity with Generative AI*, MCKINSEY & CO. (June 27, 2023), <https://www.mckinsey.com/capabilities/mckinsey-digital/our-insights/unleashing-developer-productivity-with-generative-ai> [<https://perma.cc/DN3T-782S>].

59. See, e.g., Chen, *supra* note 11.

60. See, e.g., Bamberger et al., *supra* note 13.

61. See Rajalakshmi et al., *supra* note 57 (concluding that the automated diagnostic tool's high accuracy made it appropriate for mass screenings, and noting that "availability of retina specialists . . . is a major limitation in most countries").

62. See, e.g., Megan Cerullo, *Klarna CEO Says AI Can Do the Job of 700 Workers. But Job Replacement Isn't the Biggest Issue*, CBS NEWS (Mar. 5, 2024, 3:50 PM), <https://www.cbsnews.com/news/klarna-ceo-ai-chatbot-replacing-workers-sebastian-siemiatkowski/> [<https://perma.cc/97CL-KNPZ>] (quoting the CEO of a financial technology company saying that launching an AI-assisted customer service chatbot meant that "we need the equivalent of 700 fewer full-time agents").

63. *Id.*

biased, increasing racial disparities in the medical system and causing death or serious injury.⁶⁴ Automatic pricing tools can engage in price fixing, hiking rents and charging consumers artificially high prices en masse.⁶⁵ Chatbots can deceive, including in high-stakes scenarios.⁶⁶ And while a doctor might be negligent with respect to one patient at a time, a negligently designed algorithm could affect thousands.⁶⁷

These harms are not hypothetical.⁶⁸ AI tools often exacerbate racism, sexism, and other forms of bias, making predictions and recommendations that reflect and reinforce systemic inequalities.⁶⁹ They hallucinate false information with the veneer of confidence, encouraging reliance where it is inappropriate.⁷⁰ Generative AI tools can be used to create fake pornographic images of essentially anyone, and are used in efforts to harass, intimidate, and humiliate women in particular.⁷¹ And these are just some examples of relatively direct harms of AI tools that have already begun manifesting. There are more,⁷² as well as a variety of more diffuse potential harms such

64. See, e.g., Matthew Groh et al., *Deep Learning-Aided Decision Support for Diagnosis of Skin Disease Across Skin Tones*, 30 NATURE MED. 573 (2024) (finding that a machine-learning diagnostic device improved accuracy but increased the gap in accuracy across patients of different skin tones for some physicians).

65. See, e.g., *In re RealPage, Inc., Rental Software Antitrust Litig.* (No. II), 709 F. Supp. 3d 478, 494 (M.D. Tenn. 2023) (case in which plaintiffs allege price-fixing).

66. See, e.g., Colin Lecher, *NYC's AI Chatbot Tells Businesses to Break the Law*, THE MARKUP (Mar. 29, 2024, 6:00 AM), <https://themarkup.org/news/2024/03/29/nycs-ai-chatbot-tells-businesses-to-break-the-law> [<https://perma.cc/45QE-YZQ5>] (“Five months after launch, it’s clear that while the bot appears authoritative, the information it provides on housing policy, worker rights, and rules for entrepreneurs is often incomplete and in worst-case scenarios ‘dangerously inaccurate.’”).

67. See, e.g., Jared Moore et al., *Expressing Stigma and Inappropriate Responses Prevents LLMs from Safely Replacing Mental Health Providers* 7 (unpublished manuscript), <https://arxiv.org/pdf/2504.18412> [<https://perma.cc/BXA4-5PHH>] (finding that generative AI models “do not always respond appropriately or safely” in the context of giving mental health advice, including “encouragement or facilitation of suicidal ideation”).

68. See, e.g., AI INCIDENT DATABASE, <https://incidentdatabase.ai/> [<https://perma.cc/2GUT-Y4CE>] (collecting reports of AI-related harms).

69. See, e.g., Mayson, *supra* note 40, at 2221; see also Solon Barocas & Andrew D. Selbst, *Big Data’s Disparate Impact*, 104 CALIF. L. REV. 671, 677–93 (2016) (describing how developing statistical models from large datasets, including machine learning algorithms, can often result in discriminatory outputs).

70. See, e.g., Muru Zhang, Ofir Press, William Merrill, Alisa Liu & Noah A. Smith, *How Language Model Hallucinations Can Snowball* (unpublished manuscript), <https://arxiv.org/pdf/2305.13534> [<https://perma.cc/XRG3-2BHF>].

71. See, e.g., Madyson Fitzgerald, *States Race to Restrict Deepfake Porn as It Becomes Easier to Create*, STATELINE (Apr. 10, 2024, 5:00 AM), <https://stateline.org/2024/04/10/states-race-to-restrict-deepfake-porn-as-it-becomes-easier-to-create/> [<https://perma.cc/5H8D-AW33>] (noting the prevalence of deepfake pornography and discussing instances where deepfakes have been used to target and harass women); Bobby Chesney & Danielle Citron, *Deep Fakes: A Looming Challenge for Privacy, Democracy, and National Security*, 107 CALIF. L. REV. 1753, 1773–74 (2019) (anticipating and discussing this problem).

72. See, e.g., AI INCIDENT DATABASE, *supra* note 68.

as mass job loss,⁷³ the erosion of incentives to invest in intellectual property,⁷⁴ or the degradation of the public sphere via poor-quality generated content and misinformation.⁷⁵ The fact that an individual AI tool can operate at much broader scales than an individual person thus creates both opportunities for meaningful growth and chances for highly scaled problems.⁷⁶

B. The growing world of AI legislation, regulation, and litigation

The dramatic uptick in AI use and corresponding fear of harm has led to a surge in interest in applying the tools of the law—legislation, regulation, and litigation—to this burgeoning world. In the 2024 legislative season, legislatures in more than forty states introduced AI-related legislation,⁷⁷ and one tally put the total number of proposed bills across the states and federal government at 742.⁷⁸ Legislation has been proposed to address an array of

73. See, e.g., Jill Lepore, *Are Robots Competing for Your Job?*, NEW YORKER (Feb. 25, 2019), <https://www.newyorker.com/magazine/2019/03/04/are-robots-competing-for-your-job> [https://perma.cc/7XPK-MBQK].

74. See, e.g., Frank Pasquale & Haochen Sun, Essay, *Consent and Compensation: Resolving Generative AI's Copyright Crisis*, 110 VA. L. REV. ONLINE 207, 240 (2024).

75. See, e.g., Erik Hoel, Opinion, *A.I.-Generated Garbage Is Polluting Our Culture*, N.Y. TIMES (Mar. 29, 2024), <https://www.nytimes.com/2024/03/29/opinion/ai-internet-x-youtube.html> [https://perma.cc/N6ZF-EAA3].

76. This Section has emphasized harms in the present or near term. There are also many forecasts about the impact of AI tools that predict more extreme impacts on uncertain timescales. The International Monetary Fund, for instance, recently predicted that AI “is set to profoundly change the global economy,” with “[a]lmost 40 percent of global employment . . . exposed to AI.” MAURO CAZZANIGA ET AL., INT’L MONETARY FUND, GEN-AI: ARTIFICIAL INTELLIGENCE AND THE FUTURE OF WORK 2 (2024). See generally Dylan Matthews, *How AI Could Explode the Economy and How It Could Fizzle*, VOX (Mar. 26, 2024, 4:00 AM), <https://www.vox.com/future-perfect/24108787/ai-economic-growth-explosive-automation> [https://perma.cc/MVP3-7HPA]. Similarly, there are forecasts about other extreme impacts of AI tools on society, up to and including catastrophic and existential risks. See, e.g., Yonathan Arbel, Matthew Tokson & Albert Lin, *Systemic Regulation of Artificial Intelligence*, 56 ARIZ. ST. L.J. 545, 580–83 (2024) (discussing regulatory approaches in the context of catastrophic and existential risks).

These possibilities deserve serious consideration. That is not, though the focus of this Article. There are a variety of concrete, present indicators that AI has already begun to have a significant economic role, and these are worth attending to regardless of one’s view about more extreme predictions. There is major demand for AI-related jobs, which amounted to between 1.5% and 2% of job postings in the U.S. in 2022–23. AI INDEX 2024, *supra* note 14, at 223. There is widespread adoption of AI tools across many industries. *Id.* at 258–60. And there is massive investment in AI across industries as well. *Id.* at 242–56. This Article focuses on the world of policy responses appropriate to these shorter-term, more concrete and demonstrably extant scenarios.

77. *Artificial Intelligence 2024 Legislation*, NAT’L CONF. OF STATE LEGS. (Sept. 9, 2024), <https://www.ncsl.org/technology-and-communication/artificial-intelligence-2024-legislation> [https://perma.cc/8PS8-ZK4S].

78. *Artificial Intelligence (AI) Legislation*, MULTISTATE, <https://www.multistate.ai/artificial-intelligence-ai-legislation> [https://perma.cc/ETQ3-XWTK] (identifying 742 proposed pieces of legislation, including 107 proposed at the federal level).

potential problems, ranging from tampering with elections to deepfake pornography to racial discrimination to bioterrorism.⁷⁹ On the regulatory side, a similar flurry of activity has characterized federal agencies.⁸⁰ And in the courts there has been a spate of lawsuits regarding the use of AI tools, raising issues ranging from privacy to intellectual property and heightening the odds that significant movement on AI lawmaking will come from judges.⁸¹

These efforts all confront significant challenges and questions. One question out of the starting gate is whether and when to consider disparate uses of AI tools as amenable to a single regulatory scheme, as opposed to distinct issues not connected by more than a common technological substrate.⁸² For instance, are deepfake pornography and bioterrorism connected enough to make them good targets for a single bill? Another is the extent to which new AI-specific laws are necessary, and how much we can adequately rely on existing laws that already prohibit a wide range of harmful conduct regardless of the technology used by the wrongdoer.⁸³ Will new antidiscrimination laws be necessary, for instance, or do existing laws provide an adequate basis for pursuing discrimination caused by AI systems?

These conversations make clear that much will turn on enforcement. Perhaps unsurprisingly for such a new field of regulation, many legal responses to AI contain broad standards and definitions that will be made more granular through enforcement.⁸⁴ And because of the rapid development of AI technology, a key question confronting any enforcement

79. See *Artificial Intelligence 2024 Legislation*, *supra* note 77; see also S.B. 1047, 2023–24 Leg. Sess., Reg. Sess. (Cal. 2024).

80. During the Biden Administration, for instance, the executive order on “Safe, Secure, and Trustworthy” AI required agencies to take more than 150 different actions, such as conducting studies, implementing policies, and engaging in formal rulemaking. See *Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence*, Exec. Order No. 14110, 88 Fed. Reg. 75191 (Oct. 30, 2023); see also Caroline Meinhardt et al., *By the Numbers: Tracking the AI Executive Order*, STAN. UNIV. HUM.-CENTERED A.I. (Nov. 16, 2023), <https://hai.stanford.edu/news/numbers-tracking-ai-executive-order> [<https://perma.cc/H6K5-5VQ3>].

81. See Solow-Niederman, *supra* note 2.

82. See, e.g., Tom Wheeler, *The Three Challenges of AI Regulation*, BROOKINGS (June 15, 2023), <https://www.brookings.edu/articles/the-three-challenges-of-ai-regulation/> [<https://perma.cc/E3ML-4BNM>].

83. See, e.g., Carlos Ignacio Gutierrez, *Can Existing Laws Cope with the AI Revolution?*, BROOKINGS (July 8, 2020), <https://www.brookings.edu/articles/can-existing-laws-cope-with-the-ai-revolution/> [<https://perma.cc/J6M3-QMS3>].

84. See, e.g., Colorado Artificial Intelligence Act, COLO. REV. STAT. ANN. § 6-1-1701(9)(b)(I)(A) (West 2025) (exempting AI systems from regulation under the Act if they are “intended to perform a narrow procedural task,” without further definition or elaboration).

regime is how to maintain responsiveness amid changing circumstances.⁸⁵ Given the high stakes involved in enforcement, it is unsurprising that one of the most contested parts of proposed legislation in state legislatures is whether legislation will include a private right of action, or whether enforcement will be vested exclusively in one or more government agencies.⁸⁶

But whether to allow a private right of action is only the start of the question of how to structure enforcement. In addition to considering whether to allow private enforcement in general, our legal system also must determine what kinds of private lawsuits are permitted—especially whether and when to allow aggregate litigation. And on top of that, the shape of a law’s substantive provisions may be influenced by enforcement goals as well.⁸⁷ The following sections take a deeper look at private enforcement, and at aggregate litigation in particular.

II. THE IMPORTANCE OF PRIVATE AGGREGATE ENFORCEMENT FOR AI REGULATION

Whether and how litigation is deployed will strongly influence the success of any broad effort to regulate AI tools. Litigation is part of the foundational infrastructure of the contemporary American regulatory state.⁸⁸ While typical accounts of state capacity focus on more obvious actors such as executive-branch agencies, the tools of private enforcement supplement the investigatory and enforcement activities that the state engages in to police unlawful conduct and remedy wrongs.⁸⁹ As a policy

85. See, e.g., Andrew Tutt, *An FDA for Algorithms*, 69 ADMIN. L. REV. 83, 117–18 (2017) (grounding an argument for an algorithm-focused federal agency in the “high rate of innovation in the industry” and noting the need to avoid inaction).

86. See, e.g., Brendan Bordelon, *As States Move on AI, Tech Lobbyists Are Swarming In*, POLITICO (Sept. 8, 2023, 4:23 PM), <https://www.politico.com/news/2023/09/08/tech-lobby-state-ai-efforts-00114778> [<https://perma.cc/29QK-5GMK>] (noting a “particularly” “thorny disagreement[]” over whether to include a private right of action in a prominent California proposal); Suzanne Smalley, *How Software Giant Workday Is Driving State Legislation to Regulate AI*, THE RECORD (Mar. 7, 2024), <https://therecord.media/human-resources-artificial-intelligence-state-legislation-workday> [<https://perma.cc/2KJH-K5Y7>] (noting that lobbyists for technology company Workday have succeeded in getting text from model legislation into many bills, including a lack of a private right of action); Grace Gedye & Matt Scherer, Opinion, *Are These States About to Make a Big Mistake on AI?*, POLITICO (Apr. 30, 2024, 5:00 AM), <https://www.politico.com/news/magazine/2024/04/30/ai-legislation-states-mistake-00155006> [<https://perma.cc/V8G2-NLXG>] (consumer advocates arguing that legislation should include a private right of action for the sake of strong enforcement).

87. See *infra* Part IV.

88. See SEAN FARHANG, *THE LITIGATION STATE: PUBLIC REGULATION AND PRIVATE LAWSUITS IN THE U.S.* (2010); David L. Noll & Luke P. Norris, *Federal Rules of Private Enforcement*, 108 CORNELL L. REV. 1639, 1656–63 (2023).

89. FARHANG, *supra* note 88.

choice, private enforcement has distinct advantages: It allows for decentralized enforcement of laws free from the control of state actors who may not have the resources or desire to police particular types of violations.⁹⁰ And it helps address some of the informational problems inherent in any complex society by empowering actors who are more likely to have knowledge of wrongdoing—the victims of the alleged conduct—and incentivizing them to take on the socially beneficial act of deterring future wrongdoing.⁹¹ It also serves the function of law development, especially in scenarios where novel technologies pose new legal questions.⁹² As a result, the effective deployment of private litigation will have a significant influence on the efficacy of AI-related regulations.

Within the category of private enforcement, aggregate litigation plays a distinct role whose functions are likely familiar to most lawyers.⁹³ Aggregate litigation conserves judicial resources by allowing the resolution of many claims at once. And it also allows plaintiffs to pursue claims that would not be economically viable to bring as individual cases, by bringing them together in bulk. These two related dynamics—the efficient resolution of many claims, and the creation of economic incentives to bring negative-value claims—give aggregate litigation a particularly important role as a component of public regulation.⁹⁴

These general strengths of aggregate litigation are likely to take on an outsized role in the context of litigating AI-related harms. Violations of many types of AI regulatory regimes are likely to be decentralized, implicating the strengths of private enforcement; they are also likely to involve many individuals with similar claims, implicating the efficiencies of aggregate litigation; and they are likely to involve claims that are costly to prove, suggesting that without the economies of scale of aggregate litigation, the benefits of private enforcement may be attenuated. This Part explores each of these features in turn.

90. *Id.* at 34–42 (describing different ways that enforcers can become distanced from the remedial regimes envisioned by legislators after those regimes have been enacted into legislation).

91. *Id.*; see also BRIAN T. FITZPATRICK, *THE CONSERVATIVE CASE FOR CLASS ACTIONS* (2019).

92. See Mariano-Florentino Cuéllar, Keynote, *A Common Law for the Age of Artificial Intelligence: Incremental Adjudication, Institutions, and Relational Non-Arbitrariness*, 119 COLUM. L. REV. 1773, 1780 (2019).

93. See, e.g., David Marcus, *The History of the Modern Class Action, Part I: Sturm und Drang, 1953–1980*, 90 WASH. U. L. REV. 587, 592–98 (2013) (describing the conception of the class action as having an important regulatory role).

94. *Id.*

A. Decentralization

First, many applications of artificial intelligence lend themselves to the kind of decentralized enforcement regime where private enforcement is particularly valuable. Private enforcement regimes rely on decentralized market incentives to police conduct, rather than the centralized bureaucracies of, for instance, state attorneys general or federal agencies.⁹⁵ Private parties, such as the consumers affected by a deceptive product claim or the employees affected by discriminatory actions by employers, are often better informed about wrongdoing due to their proximity to it, compared to government actors.⁹⁶ The constrained resources or incentives of government actors, meanwhile, may limit their focus to smaller numbers of violations that are particularly obvious, egregious, or easy to prove, leaving public enforcers less able to address types of misconduct that involve many different actors.⁹⁷ All else being equal, then, the more diffuse and decentralized a given category of misconduct is, the more useful the decentralized regime of private enforcement will be when it comes to policing it.

Although the commercial use of AI tools is still rapidly growing and changing, it seems likely to be an area in which private enforcement is particularly useful. Machine learning is in many ways a general purpose technology—it is a tool that can be used by many actors, in many systems, throughout the economy.⁹⁸ At least from recent experience, innovations in AI tools have diffused relatively quickly: Even though OpenAI’s GPT-4 language model represented a significant jump forward in capabilities, for instance, within a year after it was released three other companies had comparable products on the market, including a language model that Meta made freely available to encourage its widespread adoption.⁹⁹

95. See FITZPATRICK, *supra* note 91.

96. See, e.g., Stephen B. Burbank, Sean Farhang & Herbert M. Kritzer, *Private Enforcement*, 17 LEWIS & CLARK L. REV. 637, 663–64 (2013) (“Potential litigant-enforcers . . . collectively have knowledge about violations that far exceeds what the administrative state could achieve through monitoring, even under the most optimistic budget scenarios.”).

97. See, e.g., Myriam E. Gilles, *Reinventing Structural Reform Litigation: Deputizing Private Citizens in the Enforcement of Civil Rights*, 100 COLUM. L. REV. 1384, 1387 (2000) (“It was not the federal government, but private plaintiffs and community-based organizations that successfully identified and challenged unconstitutional practices in schools, prisons, mental hospitals, and other institutions.”).

98. See Avi Goldfarb, Bledi Taska & Florenta Teodoridis, *Could Machine Learning Be a General Purpose Technology? A Comparison of Emerging Technologies Using Data from Online Job Postings*, 52 RSCH. POL’Y, issue no. 1, Jan. 2023, art. no. 104653.

99. Emilia David, *Meta Says Llama 3 Beats Most Other Models, Including Gemini*, THE VERGE (Apr. 18, 2024, 2:10 PM), <https://www.theverge.com/2024/4/18/24134103/llama-3-benchmark-testing->

Similarly, the kinds of sociotechnical and market structures that are emerging around AI deployment suggest that these tools will be put to use in the hands of many different actors rather than a concentrated few. In a 2023 broad survey by McKinsey, AI tools had been adopted by companies in a wide range of industry categories, including financial services; consumer goods; healthcare systems, pharmaceuticals, and medical products; tech, media, and telecom; and business, legal, and professional services.¹⁰⁰ Many of these companies develop their own tools in house;¹⁰¹ others work with sector-specific technology firms to acquire tools with varying degrees of customization.¹⁰² And even when it comes to the kind of capital-intensive, state-of-the-art tools that raise concerns of market concentration,¹⁰³ it is unlikely that the need for regulatory scrutiny will stop at the handful of companies currently creating these tools. A current common arrangement is for these companies to make foundation models—powerful engines with important general-purpose skills like language recognition and generation—and then for other actors to fine-tune these models into specific usable products for consumers or other businesses.¹⁰⁴ As a result, responsibility and oversight for the deployment of AI tools is distributed quite broadly, even for tools that have a common origin in a small number of companies.

There is still a fair amount of uncertainty here, of course. In particular, the exact degree to which AI regulatory enforcement is centralized on a

ai-gemma-gemini-mistral [https://perma.cc/24VR-EYGC]. According to Meta, its recent open-source model, LLaMA 2, has been downloaded more than 180 million times. Mike Isaac, *How A.I. Made Mark Zuckerberg Popular Again in Silicon Valley*, N.Y. Times (May 31, 2024), https://www.nytimes.com/2024/05/29/technology/mark-zuckerberg-meta-ai.html [https://perma.cc/Y24L-LTGC].

100. AI INDEX 2024, *supra* 14, at 260 fig. 4.4.3.

101. See, e.g., Michelle Cheng, *In a Bid to Break Free from OpenAI, Companies Are Building Their Own Custom AI Chatbots*, QUARTZ (Jan. 8, 2024), https://qz.com/in-a-bid-to-break-free-from-openai-companies-are-build-1851112994 [https://perma.cc/VKN3-D6PE]; CONSUMER FIN. PROT. BUREAU 7, CHATBOTS IN CONSUMER FINANCE (2023), https://files.consumerfinance.gov/f/documents/cfpb_chatbot-issue-spotlight_2023-06.pdf [https://perma.cc/AHU3-T83C] (“As adoption of chatbots has grown, some institutions, such as Capital One, have built their own chatbot technologies by training algorithms on real customer conversations and chat logs.”).

102. See, e.g., Sara Merken, *UK Law Firm Is Latest to Partner with Legal AI Startup Harvey*, REUTERS (Sept. 21, 2023), https://www.reuters.com/legal/transactional/uk-law-firm-is-latest-partner-with-legal-ai-startup-harvey-2023-09-21/ [https://perma.cc/ZYL2-XVRR] (describing how legal technology firm Harvey creates custom large language models for individual law firms).

103. See, e.g., Jai Vipra & Anton Korinek, *Market Concentration Implications of Foundation Models: The Invisible Hand of ChatGPT*, BROOKINGS (Sept. 7, 2023), https://www.brookings.edu/articles/market-concentration-implications-of-foundation-models-the-invisible-hand-of-chatgpt/ [https://perma.cc/6G55-6B7M].

104. See, e.g., Lee et al., *supra* note 51, at 292–302; see also Olivia Moore, *How Are Consumers Using Generative AI?*, ANDREESSEN HOROWITZ (Sept. 13, 2023), https://a16z.com/how-are-consumers-using-generative-ai/ [https://perma.cc/5D3T-Y82F] (ranking fifty AI web products by monthly visits and noting that two-thirds of them are either fine-tuned versions of foundation models or a “GPT wrapper,” i.e., a user interface built on top of an existing model such as GPT-4).

small number of actors or diffuse across a broad range of actors will ultimately depend on the specific liability regime in question. One key question is where in the “pipeline” of an AI product liability attaches.¹⁰⁵ For instance, consider the following scenario: Meta trains a sophisticated model which it makes open source, releasing its weights to the world so anyone can use it. That model is then adopted and fine-tuned for more specialized medical use cases by a company that specializes in healthcare AI. That company licenses a product derived from the model to a hospital for use in clinical settings, and the hospital trains its clinicians to use it. A clinician then uses the model on a patient, and something goes wrong. Depending on the liability regime, the conduct of any of the four actors here—clinician, hospital, healthcare company, or Meta—could be the target of liability. A liability regime that tends to move liability “upstream,” toward Meta, will tend to centralize the focus of enforcement on a smaller number of actors; a liability regime that tends to move liability “downstream,” toward the clinician, will tend to decentralize the focus of enforcement and proliferate the number of regulatory targets.

At one extreme, one could imagine AI tools being like steel beams—as a practical matter, most of the liability scenarios arising from the use of steel beams are more likely to concern the manufacturer of those beams or the construction company installing those beams, and not the companies or individuals who rent space in the buildings supported by the beams.¹⁰⁶ As a tool, steel beams’ end users—the people in the buildings—have little or no exposure to liability in the typical use case. At the other extreme, one could imagine AI tools like word processors. Word processors are used all the time in liability-generating activity: writing false statements, making deceptive advertisements, and so on. But liability more or less always attaches to the end users of the tool—the people writing the document—rather than the company that made the tool on which the document was written.¹⁰⁷ Where along the spectrum AI tools end up (and there may be multiple places depending on the context and industry) will significantly influence the desirable shape of any regulatory enforcement regime. But at least so far, emerging market structures and technical practices give reason

105. See, e.g., Lemley & Casey, *supra* note 16, at 1351–78; Paul Ohm, *Focusing on Fine-Tuning: Understanding the Four Pathways for Shaping Generative AI*, 25 COLUM. SCI. & TECH. L. REV. 214, 219–31 (2024) (discussing the different stages at which generative AI models’ outputs can be influenced and considering the implications for regulatory interventions).

106. Cf. *Jimenez v. Superior Ct.*, 58 P.3d 450, 453–55 (Cal. 2002) (discussing liability in the context of manufacturing, distributing, and selling component parts of buildings).

107. Cf. *Jones v. J.B. Lippincott Co.*, 694 F. Supp. 1216, 1216–17 (D. Md. 1988) (citing *Lewin v. McCreight*, 655 F. Supp. 282 (E.D. Mich. 1987)) (publisher has no liability for the writing of a third-party author that it published, although the author themselves may be liable).

to think that enforcement regimes in a variety of key contexts will need to be able to target numerous diffuse actors rather than (or in addition to) a small number of central actors—thus increasing the relevance of private aggregate enforcement.

B. Scale

The second reason private aggregate enforcement is likely to have an important role in implementing regulations regarding machine-learning technology is because of the frequency with which AI tools are and will be deployed at scales that affect many individuals in similar ways. As discussed above, AI tools are particularly commercially appealing because of their ability to automate features of human judgment, ranging from medical diagnosis to image generation to conversational language use. This kind of automation lends itself to large-scale operations that affect many parties. The most natural way that this would happen is just a simple question of throughput: A company that is able to offer rapid, \$1 scans of patients' medical images for a basic set of diagnostics will be able to process more patients than a comparable company relying on physicians to review patient files.¹⁰⁸ And in addition to facilitating this kind of faster and cheaper version of processes that already exist, AI may also enable the creation of new kinds of processes, products, and services that affect large numbers of people—like a company that can evaluate entire databases of patient populations to look for population-health level patterns and opportunities for preventative interventions rather than focusing on individual diagnoses.¹⁰⁹

These kinds of products and services that are designed for high volumes of patients or customers will naturally give rise to scenarios where many people develop highly similar legal claims. If an investment firm's product recommendation tool turns out to run afoul of fiduciary requirements, that may taint the recommendations given to many more customers than if a single investment adviser at the firm had misunderstood their legal

108. See, e.g., Arjun Kharpal, *Zebra Medical Vision Offers \$1 AI Medical Scans on Google Cloud*, CNBC (Nov. 13, 2017, 2:19 AM), <https://www.cnbc.com/2017/11/08/google-cloud-partners-with-zebra-medical-vision-to-offer-ai-medical-image-scanning.html> [<https://perma.cc/P8U4-UTG2>].

109. See, e.g., David B. Olawade et al., *Using Artificial Intelligence to Improve Public Health: A Narrative Review*, 11 FRONTIERS PUB. HEALTH, Oct. 2023, art. no. 1196397. For an example of a commercial venture developing population health products, see also Press Release, Nanox, Nanox Announces AI Software Increases Identification of Patients with Vertebral Compression Fractures, an Early Sign of Osteoporosis, Up to Six-Fold (Mar. 12, 2024), <https://investors.nanox.vision/news-releases/news-release-details/nanox-announces-ai-software-increases-identification-patients> [<https://perma.cc/UEH9-Q7DP>].

obligations.¹¹⁰ If a medical diagnostic tool implemented across entire hospital systems is poorly designed in some way that gives rise to legal liability, the number of patients who have a claim is likely to be higher than if a single physician had a troubling pattern of misdiagnoses. In such scenarios, the typical goals of mass claiming—conserving resources, achieving uniform outcomes, and establishing global peace—are all on the table.

C. Litigation Costs

Aggregate litigation in the AI context may not simply be a way of conserving resources by taking advantage of economies of scale. It may also be necessary to ensure that some claims are even possible to bring in the first place.

Most basically, that is because there will be many claims involving allegations of algorithmic wrongdoing that are relatively small dollar value. Just as with non-AI-related claims, consumer claims involving a faulty product or service will often involve damages in the tens, hundreds, or low thousands of dollars—not enough to justify an individual suit. Some common use cases of AI fall squarely into this category,¹¹¹ and enforcement of regulations in these industries will be bolstered by private aggregate enforcement.

But a claim can be a negative-value claim not just because the amount a plaintiff can win is low; it can also be a negative-value claim if the cost to bring the claim is particularly *high*.¹¹² And litigation focused on AI tools may be particularly expensive. That is because, depending on the way that liability is established, litigating a claim involving AI wrongdoing may require a plaintiff to make complicated demonstrations about the algorithm underlying the challenged conduct.¹¹³ Making such demonstrations, in turn,

110. See, e.g., Schwarcz et al., *supra* note 49 (discussing potential legal and policy problems with automated financial advisors).

111. See, e.g., *Baker v. CVS Health Corp.*, 717 F. Supp. 3d, 188, 188 (D. Mass. 2024) (litigation over an alleged AI lie detector test involving potential statutory damages of \$500 per claimant under MASS. GEN. LAWS ANN. ch. 149, § 19B).

112. This is because a claim is “negative value” if the expected cost of pursuing it is less than the expected gains from pursuing it. So, either a lower expected gain or a higher expected cost can push the claim from being positive value to negative value.

113. There are active debates right now regarding the extent to which legal regimes should, or even can, look “inside” of an algorithmic decision as part of the process of establishing liability. See, e.g., Boris Babic & I. Glenn Cohen, *The Algorithmic Explainability “Bait and Switch”*, 108 MINN. L. REV. 857, 864 (2023). It may be the case that the difficulty in interpreting and explaining machine learning decisions ultimately counsels against liability regimes that incorporate such efforts and instead in favor of regimes like strict liability or regulatory mechanisms. See, e.g., Vladeck, *supra* note 6, at 147

could require complex discovery, sophisticated testing, and expert witnesses. Decisions made by algorithm may be difficult or impossible for humans to intuitively understand, and may need to be parsed or reformulated in ways that are more accessible—a process that itself may require computational tools, and may come with a corresponding battle of the experts over which approaches are valid or invalid.¹¹⁴ Algorithms may be sold or licensed by vendors, requiring complex third-party discovery.¹¹⁵ Plus, defendants or third parties who fear commercial or political repercussions if their algorithms are divulged may be incentivized to fight assiduously against discovery.¹¹⁶

These hurdles will all tend to increase the costs of litigation significantly. The net effect may be that many claims that would not otherwise be negative-value claims if they did not involve an algorithm will become negative-value claims simply by virtue of high litigation costs. A valid, provable claim for \$100,000 will be a positive-value claim if it costs \$30,000 to litigate, but a negative-value claim if it costs \$200,000 to litigate. If proving liability in an algorithmic context requires bringing significant technical expertise to bear regarding the particular algorithm involved, that could be enough to transform many claims from positive-value to negative-value. Aggregate litigation may therefore play a key role in maintaining the

(advocating for a strict liability regime in part because of the difficulty of establishing fault regarding driverless cars). For existing liability regimes, litigation has already developed that has involved fights over algorithms in discovery, and more seems likely as this genre of litigation becomes more common. *See, e.g.,* State Farm Fire & Cas. Co.'s Memorandum of Law in Support of Its Motion for Protective Order, *Huskey v. State Farm Fire & Cas. Co.*, No. 22-cv-07014, (N.D. Ill. Apr. 10, 2024), ECF No. 85.

114. *See, e.g.,* Babic & Cohen, *supra* note 113, at 867–70, 886–904 (describing and critiquing an approach to explaining a machine-learning decision that involves using an interpretable function to approximate an uninterpretable function, and then analyzing the behavior of the interpretable function).

115. *See, e.g.,* State Farm Fire & Casualty Co.'s Memorandum of Law in Support of Its Motion for Protective Order, *Huskey v. State Farm Fire & Cas. Co.*, No. 22-cv-07014 (N.D. Ill. Apr. 10, 2024), ECF No. 85; Opposition to State Farm's Motion for Protective Order, *Huskey v. State Farm Fire & Cas. Co.*, No. 22-cv-07014, (N.D. Ill. Apr. 25, 2024), ECF. No. 88.

116. *See, e.g.,* Flores v. Stanford, 18 Civ. 02468 (VB), 2021 WL 4441614, at *4 (S.D.N.Y. Sept. 28, 2021).

viability of private enforcement, by allowing for the amortization of expert costs across many different claims.¹¹⁷

* * *

It is important to note that private aggregate litigation is not a panacea. It will not address or solve every harm,¹¹⁸ and it can create bad incentives.¹¹⁹ Even in the context of mass disputes, there may be other systems that could be more efficient or more fair, such as the mass resolution of claims by administrative agencies,¹²⁰ or the creation of ad hoc processes or negotiations to manage specific mass problems.¹²¹ And aggregate litigation, like all litigation, is an ex post form of regulation, with all of the disadvantages that entails.¹²² The regulation of AI tools can and should rely on more forms of governance and enforcement than aggregate litigation alone.

But as this Part has argued, aggregate litigation has an important role to play. Private enforcement allows for a dispersed and decentralized response to a set of technologies that will be used extensively but in widely varied ways. And aggregation will allow both for a response to the scale of potential harms and a way to manage the costs of litigation. Especially when it comes to the kind of mass harms that have traditionally been mainstays of

117. See, e.g., Charles Silver & Lynn A. Baker, *Mass Lawsuits and the Aggregate Settlement Rule*, 32 WAKE FOREST L. REV. 733, 744 (1997) (discussing economies of scale in aggregate litigation and noting that “per capita litigation costs” can be reduced by sharing expert fees, among other things).

118. See, e.g., JULIE COHEN, KNIGHT FIRST AMEND. INST., COLUM. UNIV., HOW (NOT) TO WRITE A PRIVACY LAW 15–18 (2021), <https://s3.amazonaws.com/kfai-documents/documents/306f33954a/3.23.2021-Cohen.pdf> [<https://perma.cc/LC4L-UQZ3>].

119. See, e.g., Linda S. Mullenix, *Ending Class Actions as We Know Them: Rethinking the American Class Action*, 64 EMORY L.J. 399, 432–35 (2014).

120. See, e.g., Michael Sant’Ambrogio & Adam S. Zimmerman, *Collective Decision-Making and Administrative Justice*, in THE OXFORD HANDBOOK OF ADMINISTRATIVE JUSTICE 585 (Marc Hertogh et al., eds., 2021).

121. See, e.g., Adam S. Zimmerman, *Presidential Settlements*, 163 U. PA. L. REV. 1393, 1403–07 (2015) (describing ad hoc “presidential settlements” of major disputes); Daniel Schwarcz, *Redesigning Widespread Insurance Coverage Disputes: A Case Study of the British and American Approaches to Pandemic Business Interruption Coverage*, 71 DEPAUL L. REV. 427, 433–50 (2022) (contrasting the British and U.S. approaches to resolving mass disputes involving business interruption insurance during the COVID-19 pandemic).

122. See, e.g., Aziz Z. Huq, *Constitutional Rights in the Machine-Learning State*, 105 CORNELL L. REV. 1875, 1938–52 (2020) (discussing relative strengths and weaknesses of *ex ante* versus *ex post* approaches to regulating machine learning tools).

aggregate litigation, law enforcement in the context of AI tools will thus be more likely to be effective where private aggregate enforcement is possible.

III. THE AVAILABILITY OF AGGREGATION TO MANAGE AI HARMS

So—private aggregate enforcement has a potentially important role to play when it comes to enforcing AI-related laws. But will it be able to play that role? The last twenty years have seen significant legal changes that have made it harder to bring class actions.¹²³ Arbitration clauses with class-action waivers are now enforced;¹²⁴ defendants are more able to remove class actions from plaintiff-friendly state courts;¹²⁵ there is heightened scrutiny around Rule 23's commonality requirement;¹²⁶ increased evidentiary requirements for class certification;¹²⁷ increased scrutiny of class definitions;¹²⁸ a more demanding approach to the numerosity requirement;¹²⁹ an expansion of the adequacy requirement;¹³⁰ limitations on monetary redress in (b)(2) class actions;¹³¹ heightened scrutiny of predominance;¹³² and more.¹³³

Juxtaposed against the increasingly narrow contours of class action doctrine is the fact that AI tools will shape and structure conduct in different, and perhaps unanticipated, ways. As discussed above, AI tools allow for a novel combination of personalization and automation.¹³⁴ And while automation is likely to facilitate aggregating claims, personalization will potentially push against it. Courts, meanwhile, will face evidentiary challenges when it comes to teasing apart difficult questions of conduct

123. See generally Richard Marcus, *Bending in the Breeze: American Class Actions in the Twenty-First Century*, 65 DEPAUL L. REV. 497, 509 (2016); Myriam Gilles, *The Day Doctrine Died: Private Arbitration and the End of Law*, 2016 U. ILL. L. REV. 371, 376; Theodore Eisenberg & Kevin M. Clermont, Essay, *Plaintiphobia in the Supreme Court*, 100 CORNELL L. REV. 193 (2014); Robert H. Klonoff, *The Decline of Class Actions*, 90 WASH. U. L. REV. 729, 747–61 (2013).

124. See generally Margaret L. Moses, *Statutory Misconstruction: How the Supreme Court Created a Federal Arbitration Law Never Enacted by Congress*, 34 FLA. ST. U. L. REV. 99 (2006) (discussing the Supreme Court's increasingly broad interpretations of the FAA during the twentieth century); Gilles, *supra* note 123, at 395–401 (discussing the Supreme Court's arbitration cases involving class actions in the twenty-first century).

125. See, e.g., Helen Norton, *Reshaping Federal Jurisdiction: Congress's Latest Challenge to Judicial Review*, 41 WAKE FOREST L. REV. 1003, 1037–41 (2006).

126. See Klonoff, *supra* note 123, at 773–80 (discussing *Wal-Mart Stores v. Dukes*, 564 U.S. 338 (2011)).

127. *Id.* at 747–61.

128. *Id.* at 761–68.

129. *Id.* at 768–73.

130. *Id.* at 780–88.

131. *Id.* at 788–92.

132. *Id.* at 792–800.

133. *Id.* at 800–23.

134. See *supra* Section I.B.

involving algorithmic decision-making. This Part identifies and assesses these dynamics and their effects on the availability of class actions. Along the way, it illustrates how these dynamics are currently playing out by examining recent AI-focused litigation.

Class action doctrines are numerous and varied, and this Part does not purport to detail a comprehensive survey of, for instance, each element of class certification under Rule 23. Instead, this Part focuses on a core issue that arises in multiple important doctrinal areas: the issue of individuation, or how much it is necessary for the court to make individual determinations about class members given the conduct at issue and the relevant liability regime.¹³⁵ Individual class members might, for instance, have purchased slightly different products, or live in different locations, or been treated in slightly different ways, or be entitled to different amounts of damages. A key question becomes whether those differences defeat certification or can be accommodated within a certified class.¹³⁶ This question can arise under a number of headings, including commonality, predominance, Article III standing, or even personal jurisdiction.¹³⁷ This Part therefore focuses on this core underlying question: whether the liability regimes and conduct involved in AI-related litigation tend to require more individualized adjudicative attention, or less.

That question, in turn, depends both on the underlying facts at issue and the law governing the claims in the case. The first Section here focuses on the facts—the underlying actions enabled and incentivized by AI tools, and how they influence the availability of aggregation. The next Section turns

135. Although class actions are inherently oriented toward resolving issues for entire groups, there is some amount of adjudicative attention that must be paid to individual class members to resolve the claims in many cases, especially those brought under Rule 23(b)(3). *See, e.g.,* Allan Erbsen, *From “Predominance” to “Resolvability”: A New Approach to Regulating Class Actions*, 58 VAND. L. REV. 995, 1000 (2005) (“Courts considering whether to certify proposed class actions . . . face a recurring dilemma about how to resolve the tension between common and individual questions that arises when class members present factual circumstances that are similar, but not exactly alike.”).

136. You can think of the need for individuated attention as existing on a spectrum. *See, e.g., id.* at 1023–24 (“The significance of individualized questions in proposed class actions is a matter of degree rather than of absolutes.”). At one end of the spectrum are class actions where every element of liability can be resolved with proof that is common to the class as a whole, and the only individualized attention that may be necessary, if any, would be extremely simple proof that an individual is a member of the class in question—say, as long as a person can show a receipt that they purchased a product, they are entitled to a specific, uniform sum of money. The other end of the spectrum is the point where individual issues begin to predominate over class issues, and so the class action breaks down and is no longer certifiable. Everything in between marks the space of possible class actions—where there may be some need for individuated showings, but a class action can proceed nonetheless.

137. *See, e.g.,* David Marcus & Will Ostrander, *Class Actions, Jurisdiction, and Principle in Doctrinal Design*, 2019 BYU L. REV. 1511, 1533–44 (discussing the “injury difference problem”); *id.* at 1525–33 (discussing personal jurisdiction).

to the law, examining how the substantive rules that govern this conduct will influence when and whether aggregation can be used.

A. The influence of AI tools on conduct and aggregation

While there will obviously be a wide range of potential conduct giving rise to lawsuits, this Section identifies three relevant effects that AI tools will have in shaping conduct. First, there is a homogenizing effect, in which previously disparate actors are replaced by a single system—tending to facilitate aggregate litigation. Second, there is a differentiating effect, in which previously uniform conduct is made more individualized—tending to thwart aggregate litigation. Finally, there is a proof effect, in which the ability of class members to prove harm from AI tools will be easier or harder depending on a court’s willingness to rely on evidence derived from algorithms. Each of these effects will be important for legislators, regulators, and litigators to consider when crafting and enforcing liability regimes.

1. The homogenizing effect

First, the homogenizing effect: The use of AI tools in some contexts may make class litigation easier by replacing messy and heterogeneous decision-making processes that involve many individuals with a single algorithmic process. That, in turn, can make it easier to satisfy parts of the class action inquiry such as commonality and predominance that look to whether the plaintiff class was all treated in some common way that facilitates establishing liability with respect to the group as a whole.

Consider, for instance, *Wal-Mart v. Dukes*, in which the Supreme Court heightened the commonality inquiry for class actions.¹³⁸ In that case, a key objection that the Court majority had against certification was that the number and heterogeneity of managers responsible for hiring and promotion decisions made it difficult to establish anything that would be true across all of the decisions that affected members of the class.¹³⁹ Although the choice to allow managerial discretion was a company-wide policy, there was no “glue holding the alleged *reasons* for” challenged decisions together.¹⁴⁰

In contrast, if Wal-Mart had run its hiring decisions through a single algorithmic process, it would be more straightforward to argue that class

138. *Wal-Mart Stores v. Dukes*, 564 U.S. 338 (2011).

139. *See, e.g., id.* at 355–56 (“In such a company, demonstrating the invalidity of one manager’s use of discretion will do nothing to demonstrate the invalidity of another’s.”).

140. *Id.* at 352.

certification would be appropriate in a legal challenge to that algorithm. A single tool can be tested for bias and its design choices can be described and evaluated.¹⁴¹ A company's decision to use the tool can be assessed, and the knowledge, intention, and care behind that decision can be evaluated. Both of these inquiries may result in information that could plausibly be material to the claims of anyone affected by conduct downstream of that decision. Questions of the developer's design of the tool or the corporate customer's decision to use the tool are likely to be upstream of the later applications of that tool to individual patients, employees, consumers, etc., making it easier to argue that those questions satisfy Rule 23's commonality requirement and support the certification during the predominance inquiry as well.

There may still be hurdles, of course, such as the challenges of interpreting and evaluating algorithmic decision-making, or other features of the challenged conduct that weigh against class certification. It may be, for instance, that it is possible to prove that an AI tool that was applied to an entire class was biased, but that further proof that that bias caused harm and a particular quantum of damages requires some individual assessments of each distinct class member's circumstances. This kind of back-and-forth about common causes versus individual circumstances are the bread and butter of class litigation.¹⁴² But in general, the replacement of many decision-makers with an automated tool will tend, all else being equal, to make it more feasible to argue that the people on the receiving end of the decisions have shared a common experience, amenable to group resolution.

Perhaps unsurprisingly, early examples of litigation over this kind of activity can be found in the antitrust realm—a key place where battles are fought over what kind of coordination is and is not legally permissible. In particular, a wave of lawsuits has arisen in different contexts over pricing, in which disparate actors have allegedly ceded pricing control over to a single system. One group of these lawsuits comes from the world of real estate, with allegations that property owners and managers have unlawfully colluded by contracting with the company RealPage and outsourcing their pricing decisions to its algorithmic pricing model.¹⁴³ RealPage's software has replaced older, manual ways that landlords have set prices, including keeping records by hand and placing phone calls to competitors to check

141. See, e.g., Emily Black, John Logan Koepke, Pauline T. Kim, Solon Barocas & Mingwei Hsu, *Less Discriminatory Algorithms*, 113 GEO. L.J. 53, 96–99 (2024); Gillis, *supra* note 40, at 1247.

142. See, e.g., Erbsen, *supra* note 135, at 999–1000.

143. See, e.g., *In re RealPage, Inc., Rental Software Antitrust Litig.* (No. II), 709 F. Supp. 3d 478, 493–94 (M.D. Tenn. 2023).

prices.¹⁴⁴ RealPage's machine learning tool, in contrast, allows for the automation of management decisions for tens of millions of units across the country, based on a training dataset involving billions of price points.¹⁴⁵ Pricing decisions at this kind of scale would not be feasible for a manual, human-decision-based system. RealPage has thus taken diffuse and distinct pricing decisions made by many actors in many jurisdictions and, to at least some extent, homogenized them—running them through a single entity's software tools that output automated pricing decisions.

A similar story is playing out in the world of health insurance pricing. A spate of lawsuits has recently been filed against MultiPlan, a company that contracts with health insurers to help them set the prices at which they reimburse medical service providers.¹⁴⁶ For decades, MultiPlan assisted with pricing via negotiations with service providers.¹⁴⁷ But in recent years, MultiPlan has developed an automated price-setting tool, which is now used by an estimated 100,000 health plans covering more than sixty million people.¹⁴⁸ MultiPlan and several large health insurance companies that use its tools have now been sued by plaintiffs seeking to represent classes of medical service providers, who argue that MultiPlan's use of automated pricing across large numbers of insurers amounts to an antitrust violation.¹⁴⁹ As with the RealPage litigation, the allegations of the MultiPlan litigation depict conduct in which many disparate actors (insurance companies setting their own rates) have been substituted by an algorithmic tool that is able to benefit from and process extremely high volumes of data—hundreds of

144. See Heather Vogell, Haru Coryne & Ryan Little, *Rent Going Up? One Company's Algorithm Could Be Why.*, PROPUBLICA (Oct. 15, 2022, 5:00 AM), <https://www.propublica.org/article/yieldstar-rent-increase-realpage-rent> [<https://perma.cc/ZDF6-V9DZ>].

145. *Id.* (reporting that RealPage's software was used to manage 19.7 million units in 2020); see also *Pricing Optimization That Outperforms the Market 2%–5%*, REALPAGE <https://www.realpage.com/asset-optimization/revenue-management/> [<https://perma.cc/VPM9-S95K>] (RealPage web site claiming that their software “leverag[es] AI learnings from 20+ years and billions of transactional data points”).

146. See Alan Goforth, *MultiPlan Faces Another ‘Price Fixing’ Lawsuit, Filed by Community Health Systems*, BENEFITSPRO (May 21, 2024, 12:28 PM), <https://www.benefitspro.com/2024/05/21/community-health-systems-is-3rd-health-system-sue-multiplan-for-price-fixing/?slreturn=2024062092925> [<https://perma.cc/5GZ7-8TYQ>].

147. Chris Hamby, *In Battle Over Health Care Costs, Private Equity Plays Both Sides*, N.Y. TIMES (Apr. 7, 2024), <https://www.nytimes.com/2024/04/07/us/health-insurance-medical-bills-private-equity.html> [<https://perma.cc/M3A4-EUKR>].

148. Chris Hamby, *Insurers Reap Hidden Fees by Slashing Payments. You May Get the Bill.*, N.Y. TIMES (Apr. 9, 2024), <https://www.nytimes.com/2024/04/07/us/health-insurance-medical-bills.html> [<https://perma.cc/LLS6-R3XG>]. The exact nature of this automated tool, such as whether it has been developed via machine learning or incorporates machine learning algorithms, is unclear; MultiPlan regards details about the tool as proprietary and has sought to avoid disclosing them in litigation. See Hamby, *supra* note 148.

149. See, e.g., Class Action Complaint, *Advanced Orthopedic Ctr., Inc. v. Multiplan, Inc.*, No. 24-cv-04656, (S.D.N.Y. June 18, 2024), ECF No. 1; Class Action Complaint, *Curtis F. Robinson, M.D., Inc. v. Multiplan, Inc.*, No. 24-cv-02993 (N.D. Cal. May 17, 2024), ECF No. 1.

thousands of claims every day from more than 700 insurance company clients.¹⁵⁰

Although the volume of decisions made by the AI tools in these antitrust cases is particularly dramatic, antitrust law does not have a monopoly on the homogenizing effect of automated decision-making systems. Similar kinds of proposed class actions have arisen in recent years in the areas of consumer finance,¹⁵¹ employment law,¹⁵² and homeowners insurance,¹⁵³ just to name a few. In these scenarios, the common treatment of the proposed class by the defendant's algorithm creates grounds for plaintiffs to argue for commonality.¹⁵⁴

Procedurally, much is likely to turn in these suits on the extent to which this homogenization of decision-making can be demonstrated and shown as material to the plaintiffs' theory of liability. The defendants in the RealPage litigation, for instance, have pegged multiple lines of defense on the argument that the RealPage software provided only "recommendations," and every defendant had to make their own decision as to whether to accept or reject those recommendations.¹⁵⁵ The plaintiffs, for their part, argue that these "recommendations" were in fact mandates that were policed by RealPage in the style of a cartel, leaving little or no discretion for individual decisionmaking.¹⁵⁶ The MultiPlan litigation is less far along, but the defendants have made similar arguments to the RealPage defendants in at

150. Class Action Complaint at 52–53, *Advanced Orthopedic Ctr., Inc. v. Multiplan, Inc.*, No. 24-cv-04656 (S.D.N.Y. June 18, 2024), ECF No. 1.

151. See, e.g., Complaint, *In re Wells Fargo Mortg. Discrimination Litig.*, No. 22-cv-00990 (N.D. Cal. Feb. 17, 2022), ECF No. 1 (proposed class action involving the defendant's use of an automated underwriting system for evaluating mortgage loan applications).

152. See, e.g., *Mobley v. Workday, Inc.*, No. 23-cv-00770, 2024 WL 208529 (N.D. Cal. Jan. 19, 2024) (proposed class action involving the use of an automated job-applicant-screening tool).

153. See, e.g., *Huskey v. State Farm Fire & Cas. Co.*, No. 22-cv-07014, 2023 WL 5848164 (N.D. Ill. Sep. 11, 2023) (proposed class action involving the use by an insurance company of an automated fraud-detection tool).

154. See, e.g., Plaintiffs' Notice of Motion and Motion for Class Certification at 14, *In re Wells Fargo Mortg. Discrimination Litig.*, No. 22-cv-00990 (N.D. Cal. Apr. 25, 2024), ECF No. 204 (arguing that a claim involving the defendant's use of an automated underwriting system for evaluating mortgage loan applications "presents a straightforward case for class certification").

155. See, e.g., Defendants' Reply Memorandum of Law in Support of Motion to Dismiss Multifamily Plaintiffs' Second Amended Consolidated Class Action Complaint at 2, *In re RealPage, Inc., Rental Software Antitrust Litig.* (No. II), 709 F. Supp. 3d 478 (M.D. Tenn. 2023) (No. 23-md-03071), ECF No. 639 (arguing that lumping defendants together without accounting for their different responses to RealPage's recommendations is impermissible "group pleading"); *id.* at 6 (arguing that the fact that defendants rejected some recommended prices defeats allegations of parallel conduct); *id.* at 14 (saying that the possibility that defendants rejected price recommendations contributes to the argument that plaintiffs lack antitrust standing).

156. See, e.g., Plaintiffs' Memorandum of Law in Opposition to Defendants' Motion to Dismiss the Multifamily Consolidated Amended Complaint at 3–4, *In re RealPage, Inc., Rental Software Antitrust Litig.* (No. II), 709 F. Supp. 3d 478 (M.D. Tenn. 2023) (No. 23-md-03071), ECF No. 623.

least one case.¹⁵⁷ Ultimately, both substantive issues of antitrust liability and procedural issues of class certification will likely hinge on the extent to which plaintiffs can show that these algorithmic tools created a single course of conduct where previously there had been many distinct courses. These showings are likely to depend on both the state of the law as well as the state of our technical understanding of AI tools at any given point in time.¹⁵⁸ And this kind of fight—over how much individualized inquiry is necessary when a class files a lawsuit about an algorithm’s recommendations—is likely to come up time and time again in the years ahead.

2. *The differentiating effect*

Not all uses of algorithms will tend toward this kind of concentrated conduct that favors class certification. In addition to a homogenizing effect, there is also likely to be a differentiating effect that will push against class actions. In some circumstances, AI tools will substitute not for heterogeneous decision-making, but instead for a uniform policy or process. AI tools allow for more nuanced and personalized automation than previous kinds of software.¹⁵⁹ This will enable some actors to replace uniform approaches with tools that tolerate more variable outcomes. An investment brokerage, for instance, may be able to offer an increased array of options mediated by a chatbot, rather than a smaller number of choices mediated by staff or just selected by individual customers after reading summaries.¹⁶⁰

That arrangement, in turn, would likely lead to more heterogeneous treatment of customers in ways that militate against class certification. This heterogeneous treatment could take a couple of forms. First, because the products on offer can be more varied and complex, there may be more

157. See Multiplan, Inc.’s Memorandum of Law in Support of Its Motion to Dismiss the Complaint Pursuant to Federal Rule of Civil Procedure 12(B)(6) at 4, 10, *Adventist Health Sys. Sunbelt Healthcare Corp. v. Multiplan, Inc.*, No. 23-cv-07031 (S.D.N.Y. Dec. 12, 2023) (arguing that because price proposals are “just the opening volley” and that MultiPlan’s individual clients “decide how they use” its tools, its practices do not restrain trade).

158. Technical evidence about what is going on within AI systems is likely to be an important tool in at least some cases for determining whether potential class members’ claims are similar enough for class certification. For instance, recent evidence indicates that within a large language model, different books within the training data are “memorized” by the model to different degrees—with potential implications for class actions based around intellectual property claims. See A. Feder Cooper et al., *Extracting Memorized Pieces of (Copyrighted) Books from Open-Weight Language Models* 7 (unpublished manuscript), <https://arxiv.org/pdf/2505.12546> [<https://perma.cc/HX38-2V3V>].

159. See, e.g., CITI, *AI IN FINANCE: BOT, BANK & BEYOND* (2024), https://ir.citi.com/gps/9j79xHlavfPi785TYiSciff00j4I0D52fi9LrahsLZEo6MpT4aM7SpwSFagAL9CIukqn2fwiJ_GNvDsLy4b6XEjftdK1abu [<https://perma.cc/32BP-NK66>] (discussing the “hyper personalized” products and services that AI tools are likely to enable in the financial sector).

160. *Id.*

material differences between the various products that customers select, and there may be fewer individuals who have purchased any given product. That may make it less feasible to bring a class action regarding any particular product, and also make it difficult to certify a class across products. If a company is truly able to use AI to make “products for a market of one,” in other words, it might be harder to bring product liability actions as a broad class.¹⁶¹

Second, the interactions between company and customer might get more varied. In the investment brokerage example, the chatbot’s conversations with individual customers might be highly personalized and path-dependent, such that no customer has the exact same conversation with the chatbot. For liability regimes in which a defendant’s representations of information feature heavily, such as the kind of deception claims common in consumer protection cases, the need to parse individual representations from a company’s chatbot may make class certification infeasible.

Particularly relevant here is the fact that some AI systems, and many generative AI systems in particular, are non-deterministic.¹⁶² In the context of machine learning models, a system is non-deterministic if it can be given the same inputs but produce different outputs.¹⁶³ Some non-determinism is a difficult-to-avoid feature of the process of training machine learning models,¹⁶⁴ but for large language models in particular developers intentionally add even more variation into the model’s response to achieve what are seen as better results.¹⁶⁵ When it comes to aggregate litigation, though, such non-deterministic behavior can be a hurdle to satisfying commonality and predominance. Non-deterministic outputs mean that it’s possible for similarly situated people to be treated differently by an AI

161. *Id.*

162. See, e.g., A. Feder Cooper, Jonathan Frankle & Christopher De Sa, *Non-Determinism and the Lawlessness of Machine Learning Code*, 2022 CSLAW 1, <https://arxiv.org/pdf/2206.11834> [<https://perma.cc/PE77-G8DB>]. Non-determinism can occur within a single model (i.e., a single model can give different outputs from the same input), or it can occur as a property across models (i.e., even if multiple models are trained on the same data, they may produce different outputs from the same input). *Id.* at 3–6. And while some non-determinism is “stochastic,” meaning that it can be described using familiar probabilistic reasoning, other non-determinism is not stochastic and does not follow the kinds of patterns that can be analyzed with probabilities. *Id.* at 3.

163. *Id.* at 2.

164. *Id.* at 3–4.

165. See, e.g., Stephen Wolfram, *What Is ChatGPT Doing . . . and Why Does It Work?*, STEPHEN WOLFRAM WRITINGS (Feb. 14, 2023), <https://writings.stephenwolfram.com/2023/02/what-is-chatgpt-doing-and-why-does-it-work/> [<https://perma.cc/4Q3P-7UBY>] (describing how changes to the “temperature” setting can qualitatively improve the results of language models).

tool.¹⁶⁶ In the investment brokerage example, two customers with similar financial situations and preferences may be steered in different directions.¹⁶⁷ If those different directions are material to the theory of liability at issue in the case, it would be more difficult to argue that commonality and predominance are satisfied even if the customers all interact with the same model.

It's important not to overstate the force of this effect. While many AI use cases will involve a differentiated and personalized product or service, the same is true of many of the existing products or services that the tools will augment or replace. An AI chatbot that has a "conversation" with a user and offers them a slate of investment products and advice may not make it any more challenging to certify a class than in a context where that conversation and advice happens with a human. It is only where there is a marginal increase in the differentiated conduct existing in the world, and where that conduct is relevant to a plaintiff's theory of liability, that there will be a new hurdle for class certification. Nonetheless, it does seem likely that there will be both newly differentiated goods and services and entirely new goods and services,¹⁶⁸ making this differentiating effect an important facet of the burgeoning AI litigation landscape.

3. *The proof effect*

The homogenizing and differentiating effects just discussed are effects that AI tools will have on the underlying facts that give rise to the disputes brought into court. But AI tools will also influence how courts are able to process those facts. Some of these changes will be upgrades to relatively simple administrative functions, such as communicating with class members, processing claims forms, and so on.¹⁶⁹ But some changes will be

166. See, e.g., Lecher, *supra* note 66 ("[T]he bot doesn't return the same responses to queries every time. At one point, it told a Markup reporter that landlords *did* have to accept housing vouchers, but when ten separate Markup staffers asked the same question, the bot told all of them no, buildings did not have to accept housing vouchers.").

167. A distinct problem raised by non-deterministic behavior is that it can make it difficult to establish causation at the individual level. See *infra* notes 214–234 and accompanying text.

168. See, e.g., AI IN FINANCE, *supra* note 159.

169. See, e.g., *Document Automation*, LEGALTECH HUB, <https://www.legaltechnologyhub.com/topics/contracts/document-automation/> [<https://perma.cc/ZJD5-BYNY>] (listing 137 legal technology tools designed to assist with document automation); *Receptionist Service for Law Firms*, SMITH.AI, <https://smith.ai/industries/legal-law-firms-answering-service> [<https://perma.cc/546W-KLSH>]; see also *Stay Client-Focused During Seasonal Rush with Call Answering*, SMITH.AI, <https://smith.ai/case-studies/youngblood-group> [<https://perma.cc/4W3V-JC47>] (noting the capacity to answer routine questions and book meetings); Susan Raridon Lambreth & Amanda C. Chaboryk, *Generative AI: A New Tool for Legal Project Management*, LAWVISION (Jan. 9, 2024), <https://lawvision.com/generative-ai-a-new-tool-for-legal-project-management/> [<https://perma.cc/6P3V-X7BU>].

more directly connected to core adjudicative functions. In particular, the perennial challenge of balancing the benefits of group claim resolution with the necessities of distinguishing between claimants may look different where AI tools can provide low-cost ways of assessing individual circumstances. And where it is easier to manage these kinds of individual assessments, it will be easier to certify a class. There will therefore be a kind of “proof effect” for class actions, in which the cost-benefit ratio of certifying a class will change depending on how the use of AI tools affects courts’ ability to make individualized proof inquiries.

Consider, for instance, the following scenario:

Algorithmic Diagnosis: A medical services provider sells a diagnostic tool to screen for a variety of early-stage cancers.¹⁷⁰ The tool is powered by a machine learning algorithm called Algorithm A. Many people are screened by Algorithm A; as with any diagnostic tool, there are some false positives and some false negatives. But it later turns out that Algorithm A was negligently designed, and did not meet the relevant standard of care. An alternative machine learning algorithm, Algorithm Z, meets the standard of care, and it is possible to take a copy of a person’s data that was run through Algorithm A and run it through Algorithm Z to see what the results would have been had the better algorithm been used. Under the governing liability regime, a person will have a valid claim against the medical service provider if (1) they were misdiagnosed by Algorithm A, (2) they would have been accurately diagnosed by Algorithm Z, and (3) they suffered an injury as a result of not being accurately diagnosed.

This type of claim is in some sense a garden-variety medical malpractice claim: The plaintiff was not treated with the standard of care, and if they had been treated with the standard of care they would have avoided some injury.¹⁷¹ The involvement of an algorithm does not necessarily have to change that,¹⁷² but it does affect the way that the claim is established. A

170. See, e.g., Marlene Cimon, *A Blood Test to Detect Cancer? Some Patients Are Using Them Already.*, WASH. POST (Apr. 16, 2024), <https://www.washingtonpost.com/wellness/interactive/2024/cancer-blood-test-screening/> [<https://perma.cc/3K74-6NSK>].

171. See, e.g., *Gardner v. Pawliw*, 696 A.2d 599, 608 (N.J. 1997).

172. It is possible that the use of an algorithmic diagnostic device could entail a different liability regime, such as products liability rather than medical malpractice. The specific liability regime is not the focus of this example—the point is just that there may be *some* remedial regimes that mean it will be useful, or possibly necessary, to “apply” an algorithm in the course of establishing liability in some circumstances.

similar case without machine learning diagnostics would typically involve competing expert witnesses testifying about the diagnostic reasoning used by the defendant, and why it was or was not adequate under the standard of care.¹⁷³ With the machine learning algorithm added, there would likely still be expert witnesses, but this time they will be focused on establishing whether Algorithm A met the standard of care and whether Algorithm Z can be used to establish that the claimant could have been accurately diagnosed.

In a case involving just one person, it seems like litigation in the case involving an algorithm is just more complex—there is the same battle of the experts, with the added complication of an opaque machine learning tool. In this way, the involvement of algorithms may make many types of cases more difficult to litigate.

But where there are many plaintiffs, the role of algorithms opens up a new possibility: that the claims can proceed as a class. Rather than requiring individual expert evaluations of each plaintiff's medical data, it would be possible to establish that, in general, Algorithm Z can provide the necessary comparison with Algorithm A to demonstrate that a person was misdiagnosed in a way that could have been corrected by a standard-of-care tool. Once that proposition is established as a general matter, it would be possible in principle to process the class members through Algorithm Z without the same need for individualized expert assessment. There may still be a need for some individual assessments, such as a damages inquiry, but the ability to substitute an algorithm for expert assessment would be a significant step in favor of class resolution.

This scenario is plausibly generalizable to a range of potential litigation. A brokerage's robo-advisor powered by a machine-learning tool may give self-dealing advice, and it may take another machine-learning tool to tease out who received an inappropriate recommendation. Or a health insurer may use a machine-learning-based tool for its reimbursement practices in a way that amounts to an antitrust violation,¹⁷⁴ and it may take a more "neutral" tool to establish which payments' prices were affected. As the complexity of machine-learning tools increases, and as more commercial and governmental practices employ these tools, the number of scenarios is likely to grow in which it is necessary to turn to a machine-learning tool to assess

173. Cf. *Kipfinger v. Great Falls Obstetrical & Gynecological Assocs.*, 525 P.3d 1183, 1196–97 (Mont. 2023) (discussing expert testimony in a medical malpractice case).

174. See, e.g., Chris Hamby, *Collusion in Health Care Pricing? Regulators Are Asked to Investigate*, N.Y. TIMES (May 1, 2024), <https://www.nytimes.com/2024/05/01/us/multiplan-health-insurance-price-fixing.html> [<https://perma.cc/5V8X-7ZYU>].

and evaluate another tool.¹⁷⁵ And when those scenarios give rise to potential liability, courts will need to be able to accommodate evidence that incorporates the use of machine learning to resolve disputes. AI tools may thus be useful, or at times even necessary, for fact-finding efforts in certain kinds of complex disputes.

But the fact that such cases may exist, and that the class device may be useful for resolving them, does not mean that class action doctrine will allow these cases to proceed as a class. Class action doctrine has had a wary relationship to statistical evidence, perhaps most emblematically in Justice Scalia's pejorative characterization of using statistical sampling as conducting "Trial by Formula" in the *Wal-Mart v. Dukes* majority opinion.¹⁷⁶

That is not to say that statistical evidence is forbidden; just a few years later, the Supreme Court affirmed that plaintiff classes can use statistical sampling to establish liability where any class member suing individually would be permitted to use such evidence to establish liability.¹⁷⁷ In a scenario like *Algorithmic Diagnosis*, the use of a machine-learning tool would likely be a reasonable way to establish liability in an individual case. But machine learning tools are fundamentally a form of statistical reasoning, involving the use of a large sample of data on which a model is trained to make an inference about a particular input. And as Robert Bone has described, the use of "statistical techniques to adjudicate large case aggregations[] is highly controversial," even though there are strong normative arguments in favor of using statistical reasoning and evidence in aggregate litigation.¹⁷⁸ The successful use of litigation to resolve disputes of this type thus will depend on courts' ability and willingness to adduce algorithmic evidence, as well as the type of liability regime involved.

175. Work on interpreting neural networks, for instance, has used algorithms designed specifically to facilitate interpretation to help understand the inner workings of other neural networks. See, e.g., Trenton Bricken et al., *Towards Monosemanticity: Decomposing Language Models with Dictionary Learning*, ANTHROPIC (Oct. 4, 2023), <https://transformer-circuits.pub/2023/monosemantic-features> [https://perma.cc/6L9A-MXQ8].

176. See *Wal-Mart Stores, Inc. v. Dukes*, 564 U.S. 338, 367 (2011).

177. *Tyson Foods, Inc. v. Bouaphakeo*, 577 U.S. 442, 455 (2016); see also Hillel J. Bavli & John Kenneth Felter, *The Admissibility of Sampling Evidence to Prove Individual Damages in Class Actions*, 59 B.C. L. REV. 655, 659–69 (2018) (discussing the different treatments of statistical evidence in *Wal-Mart v. Dukes* and *Tyson Foods*).

178. See Robert G. Bone, *Tyson Foods and the Future of Statistical Adjudication*, 95 N.C. L. REV. 607, 609, 655–70 (2017).

* * *

The adoption of AI tools is likely to influence the certifiability of actions challenging a wide range of conduct. But whether that influence pushes in favor or against certification depends on whether AI is adopted to reduce the heterogeneity within a given process or increase it. Many uses of AI tools will be more complicated than the examples just given—for instance, many use cases are likely to involve AI tools and recommendations augmenting human effort, rather than simply substituting wholesale for human judgments.¹⁷⁹ But given the ability of AI tools to facilitate more nuanced automated processing at large scales, the increasing adoption of those tools by commercial and government enterprises is likely to have repercussions for the law’s mechanisms for handling aggregated disputes.

B. Responses by liability regimes that facilitate aggregation

The previous Section addressed how AI tools will shape conduct in ways that make aggregate enforcement more or less feasible. But the availability of aggregate enforcement also depends on what liability rules are in play.¹⁸⁰ The fact that a group of claimants was treated in some similar way might not matter much if that similar treatment is not material to their theory of liability, or if other material facts are highly dissimilar throughout the proposed class.

As a result, effective governance of AI tools will depend on how the substantive regimes that we arrive at facilitate or thwart enforcement procedures. And although we are still in the early days of AI regulation and litigation, there are already some lessons to be learned. This Section identifies two of those lessons—how expanding liability “upstream” can facilitate class actions and how deemphasizing causation can facilitate class actions, both of which can be done in plausible ways in the AI context.

1. Moving liability upstream

Although there may be many new legal regimes built to handle the explosion of AI tools, litigants and government agencies are also relying heavily on existing liability frameworks. Decades-old laws regarding antidiscrimination, consumer protection, and more have been the basis of

179. See, e.g., Price et al., in RESEARCH HANDBOOK ON HEALTH, AI, AND THE LAW, *supra* note 6, at 150, 151–55 (discussing physician liability in the context of AI recommendations).

180. See, e.g., Erbsen, *supra* note 135, at 999 n.3 (“The nature and significance of individual issues is a function of the substantive liability and damage rules applicable to asserted claims and defenses.”).

numerous lawsuits, enforcement actions, and regulatory warnings.¹⁸¹ But one of the concerns raised regarding these laws is that they might not reach the right actors.¹⁸² Laws intended to regulate conduct by a decision-maker like a manager or a physician, for instance, may not be as effective at reaching the designers of AI tools that now assist in that decision-making, even though those tools may raise similar concerns of bias or safety.¹⁸³

One approach to this problem has been to widen the parameters of existing liability regimes so that they capture these actors. An early example of this path is Colorado’s recently passed AI legislation, generally regarded as the first state attempt at a broad freestanding AI bill.¹⁸⁴ The bill focuses on antidiscrimination, and incorporates existing antidiscrimination law.¹⁸⁵ Its key provisions define “algorithmic discrimination” to include “any condition in which the use of an artificial intelligence system results in an unlawful differential treatment or impact” disfavoring a legally protected group,¹⁸⁶ and then requires the developers of certain AI systems to “use reasonable care” to protect against the risk of algorithmic discrimination resulting from those systems.¹⁸⁷ In other words, the law piggybacks on existing antidiscrimination laws, and says to developers: take care to make sure that your products don’t result in these laws being violated.

This legislation was likely not developed with aggregate litigation in mind—in fact, Colorado’s bill forecloses private enforcement, explicitly limiting enforcement of the bill’s provisions to the state’s Attorney General.¹⁸⁸ But this form of extending liability is not limited to Colorado’s

181. See, e.g., *Huskey v. State Farm Fire & Cas. Co.*, No. 22-cv-7014, 2023 WL 5848164 (N.D. Ill. Sept. 11, 2023); *Consumer Financial Protection Circular 2022-03: Adverse Action Notification Requirements in Connection with Credit Decisions Based on Complex Algorithms*, CONSUMER FIN. PROT. BUREAU (May 26, 2022), <https://www.consumerfinance.gov/compliance/circulars/circular-2022-03-adverse-action-notification-requirements-in-connection-with-credit-decisions-based-on-complex-algorithms/> [<https://perma.cc/6J5M-9FWK>]; *In re Everalbum, Inc.*, No. 1923172, 2021 WL 2012450 (F.T.C. May 6, 2021), https://www.ftc.gov/system/files/documents/cases/1923172_-_everalbum_decision_final.pdf [<https://perma.cc/AB72-PCQ3>].

182. See, e.g., Charlotte A. Tschider, *Humans Outside the Loop*, 26 YALE J.L. & TECH. 324, 355–97 (2024).

183. *Id.*; see also *Mobley v. Workday, Inc.*, No. 23-cv-00770, 2024 WL 208529, at *5–6 (N.D. Cal. Jan. 19, 2024) (dismissing a complaint against a software company because the complaint’s allegations were insufficient to demonstrate that the company was an “employment agency” under Title VII).

184. See Jason I. Epstein et al., *Understanding the CAIA: Colorado’s Groundbreaking Approach to AI Regulation*, NAT’L L. REV. (June 20, 2024), <https://natlawreview.com/article/understanding-caia-colorados-groundbreaking-approach-ai-regulation> [<https://perma.cc/UAR4-X3N7>]; S.B. 24-205, 74th Gen. Assemb., 2d. Reg. Sess. (Colo. 2024) (enacted May 17, 2024), https://leg.colorado.gov/sites/default/files/2024a_205_signed.pdf [<https://perma.cc/5AXG-QG3G>].

185. See COLO. REV. STAT. ANN. § 6-1-1701 (West 2025).

186. § 6-1-1701.

187. § 6-1-1702(1).

188. § 6-1-1706.

legislation, although that is a particularly clear example of it. The FTC, for instance, has also sought comment on a proposal to extend liability for AI products used to impersonate businesses and government officials.¹⁸⁹ That impersonation is already unlawful, but the FTC's proposal would move liability upstream to producers of goods or services (such as AI tools) if those producers "know or have reason to know" that their products will be used for this kind of unlawful impersonation.¹⁹⁰ And beyond agency regulations, private parties can also seek via litigation to have existing law interpreted in a way that attaches liability to a course of conduct further "upstream" than it might otherwise go.

Where private lawsuits are allowed, this adjustment to liability will make aggregate litigation easier. Liability in these circumstances can be understood as moving upstream because it moves liability for a good or service further back toward that good or service's creation, and away from the "downstream" actions of individual end users. That matters because, as discussed earlier, those downstream actions by individuals are more likely to be heterogeneous and raise issues that require individualized determination.¹⁹¹ Upstream conduct, in contrast, is more likely to be conduct that is materially identical, or at least more similar, with respect to everyone who is injured by that conduct.

Take, for instance, the current spate of intellectual property lawsuits against generative AI companies.¹⁹² Although the cases present a variety of theories,¹⁹³ two main alleged copyright violations occur at different points in the generative AI supply chain.¹⁹⁴ First, one type of alleged copyright violation occurs "upstream" during the course of training machine-learning models on copyrighted information: for instance, the allegation that OpenAI violated the Copyright Act when it trained its language models on the New York Times's copyrighted articles.¹⁹⁵ Second, a different type of alleged copyright violation occurs more "downstream" when trained models generate outputs that infringe on copyrights: for instance, the allegation that OpenAI violates the Copyright Act each time its models produce copies or derivatives of the New York Times's copyrighted articles.¹⁹⁶

189See Trade Regulation Rule on Impersonation of Government and Businesses, 89 Fed. Reg. 15072 (proposed Mar. 1, 2024) (to be codified at 16 C.F.R. pt. 461).

190. *Id.* at 15079.

191. See *supra* Section II.A.

192. See, e.g., Solow-Niederman, *supra* note 2, at 273–74 (describing several suits bringing a range of theories).

193. *Id.*

194. See Lee et al., *supra* note 51, at 344–45 (discussing the generative AI supply chain).

195. See Complaint at 60–62, *N.Y. Times v. Microsoft Corp.*, No. 23-cv-11195 (S.D.N.Y. Dec. 27, 2023), ECF No. 1.

196. *Id.* at 163.

Regardless of the strength of those allegations on the merits of the copyright issue, the first type of upstream allegation is much stronger as a potential basis for class certification. That is because the factual lynchpin for members of a proposed class of IP holders would be whether their copyrighted material was contained in the database on which the defendant trained its models.¹⁹⁷ This is an individualized question—different class members would have different IP that needs to be found in the training data—but it may be relatively tractable, depending on the records that exist. In a variety of circumstances, such as where a defendant retains a searchable copy of their training data, or records indicate that their model was trained on a known corpus of materials that is searchable or well indexed, it will be possible to establish that an IP holder’s material was trained on, and doing so may be very straightforward.¹⁹⁸

In contrast, the second, downstream theory—which pegs copyright violations to the models’ outputs, rather than their training—would be more difficult to use as the basis for a certified class. That is true for several reasons. Some of those reasons are technical: Certain works in a model’s training data may be more likely to be output in identical or nearly identical form than other works, based on factors like the simplicity of the work and the frequency of its appearance in the training data.¹⁹⁹ Other reasons have more to do with the way liability can be shown. Whether a given output is a close enough copy of a copyrighted work to constitute infringement, for instance, may need to be a case-by-case determination that would be infeasible to do at scale.²⁰⁰ Or the role of the individual user may matter; for instance, liability in the copyright context could plausibly depend in part on whether the user prompted the model in a way intentionally designed to circumvent safeguards that were put in place to avoid generating

197. See, e.g., Lee et al., *supra* note 51, at 344–45.

198. See, e.g., *id.* at 345–46. This is not to suggest at all that every class action alleging this type of theory would be simple. There may be many circumstances where it is difficult or impossible to establish conclusively that a given item was part of a training set. *Id.* However, it is also possible to imagine classes constructed in ways designed specifically for scenarios where proving class membership and liability (under this theory) is straightforward, such as a class defined as all IP holders whose content is in a given corpus, in a lawsuit against a defendant who has admitted to training on that corpus. The Common Crawl, for instance, is a corpus that is popular for training language models that is also publicly available, making it relatively easy for content owners to know if their material is in the corpus. See, e.g., Kate Knibbs, *Publishers Target Common Crawl in Fight Over AI Training Data*, WIRED (June 13, 2024, 11:21 AM), <https://www.wired.com/story/the-fight-against-ai-comes-to-a-foundational-data-set/> [<https://perma.cc/N39R-V359>].

199. See Matthew Sag, *Fairness and Fair Use in Generative AI*, 92 FORDHAM L. REV. 1887, 1911–12 (2024).

200. See, e.g., Lee et al., *supra* note 51, at 330–31, 337–44.

copyrighted material.²⁰¹ Where these types of individual determinations are fact-intensive, they may functionally thwart the ability to certify a class.

Although copyright has a variety of distinctive doctrines, the increased ease of certification when challenging upstream conduct generalizes to other domains as well. As discussed above, one of the reasons that AI tools' "homogenizing effect" is noteworthy is because it is easier to certify a class based on a single course of conduct that materially affected class members in an identical or near-identical way.²⁰² Actions closer to a user's individual interaction with a good or service are more likely to be heterogeneous, whether that manifests as a generative AI chatbot's different statements made to many different customers or as many applicants' different responses to screening questions in a job interview. Actions further upstream from that conduct—the training of a model for a generative AI system or an automated hiring program—are more likely to be homogeneous with respect to their effect on members of the class. Allowing liability to attach to that more upstream conduct thus facilitates class certification, in addition to whatever other advantages or disadvantages it might have.

2. *Minimizing causation requirements*

Another way that liability regimes can facilitate class actions is by minimizing the causal proofs that individuals need to make to establish liability. The most straightforward way of doing this is via substantive rules that can be thought of as "per se liability" regimes. In these scenarios, the application of a particular algorithmic tool to a class member is itself a sufficient basis for liability, based on the inherent features of the tool in question. No further proof is needed. Most importantly, for purposes of differentiating this type of suit from other types of suits alleging algorithmic wrongdoing, liability does not rest on the outcome of an algorithmic decision. As a result, there need be no individual inquiry into what the outcome of the algorithmic process was, and/or whether there "should have" been some other type of outcome. An individual need only establish that they were subjected to some sort of treatment by an algorithmic tool that was *per se* impermissible.

One example of this type of lawsuit is *Baker v. CVS*, a prospective class action lawsuit pending in federal court in Massachusetts.²⁰³ In *Baker*, the

201. See, e.g., *id.* at 350–53 (discussing the concept of volition in cases of direct infringement).

202. See *supra* Section III.A.

203. See *Baker v. CVS Health Corp.*, 717 F. Supp. 3d 188 (D. Mass 2024).

plaintiff applied for a job with the CVS chain of pharmacy stores.²⁰⁴ Part of the application process involved a video interview developed by another company, HireVue, in which applicants were asked to answer questions on video and HireVue then used an AI tool to evaluate the applicants—allegedly including an assessment as to the applicants’ honesty during the interview.²⁰⁵ The named plaintiff in *Baker* is an applicant who says that this video assessment amounts to a lie detector test that violates Massachusetts’ statutory prohibition on conditioning employment on a lie detector test.²⁰⁶ Under the plaintiff’s theory, CVS would be liable for a statutory violation to anyone to whom the test was administered, with no need to inquire as to the results of the algorithmic process for each individual or to otherwise scrutinize the decision-making of the AI tool beyond establishing that it did, in fact, qualify as a prohibited lie detector under Massachusetts law.²⁰⁷

More examples of this category come from a variety of suits that allege algorithmic misconduct in violation of biometric privacy protection statutes. In this type of lawsuit, a plaintiff alleges that a company deployed an AI tool to interact with the plaintiff, and in the course of interacting with the customer the tool stored the plaintiff’s biometric information in violation of a law such as Illinois’s Biometric Information Privacy Act, or BIPA.²⁰⁸ The defendant’s alleged liability arises not from some flaw in the AI tool’s decision-making, or the gathering of data to make an AI tool, but rather from the use of an AI tool that is designed to gather impermissible data in the course of its operations.²⁰⁹ As with *Baker*, liability under these theories is established without needing to investigate the individual validity of any algorithmic decision; instead, the challenged conduct is a uniform action taken by the algorithm with respect to all people that it interacts with.

204. *Id.* at 189.

205. *Id.* at 191.

206. *Id.* at 189.

207. *Id.* This case raises an important Article III question, which will likely be raised by other cases in the category of “per se liability” classes as well. That question is whether a plaintiff has standing to pursue a violation of a statute per se, or whether Article III standing’s injury requirement forces them to show some sort of harm arising as a consequence of the statutory violation. *Id.* The outcome of that question will likely hinge on the right involved in any particular case. In *Baker*, for instance, the court held that the violation of the statutory right to notice was an informational injury sufficient for standing; but it noted that that injury might not be sufficient for applicants who had not in fact been subjected to the lie detector test. *Id.* at 192 (discussing *TransUnion LLC v. Ramirez*, 594 U.S. 413 (2021)).

208. *See, e.g.*, Class Action Complaint at 1, *Garcia v. Domino’s Pizza, Inc.*, No. 24-cv-02090 (N.D. Ill. Mar. 13, 2024), ECF No. 1; *Delgado v. Meta Platforms, Inc.*, 718 F. Supp. 3d 1146, 1151 (N.D. Cal. 2024).

209. *See, e.g.*, Class Action Complaint at 1, *Garcia v. Domino’s Pizza, Inc.*, No. 24-cv-02090 (N.D. Ill. Mar. 13, 2024), ECF No. 1; *Delgado v. Meta Platforms, Inc.*, 718 F. Supp. 3d 1146, 1151 (N.D. Cal. 2024).

It is easy to imagine a role for this type of liability regime in the context of other AI harms as well. Concerns about AI's role in targeting individuals for political misinformation campaigns, for instance, may militate in favor of a *per se* liability regime where an individual need show only that they were targeted rather than having to also prove any sort of economic or emotional harm from that targeting.²¹⁰ Or, in an analogy to the *per se* version of the tort of defamation, one plausible liability regime for the creation of fake nonconsensual intimate imagery would be to hold the creator or disseminator of such imagery liable for its creation *per se*, without a need for the victim to prove specific follow-on consequences like reputational harm.²¹¹ There may be strong normative grounds for a kind of *per se* liability regime in scenarios like these where the harm that policymakers are concerned about is able to be presumed from the conduct itself, and/or where proving harm would be difficult enough that policymakers would be concerned about chilling the deterrent or compensatory powers of private enforcement.

This type of liability regime has the potential to make class litigation much easier. As long as the legal theory is sound, and liability truly does arise uniformly and simply from the interaction of each class member with the algorithm, then the doctrines surrounding class actions should not form a particularly significant hurdle. There will, of course, be all of the usual limitations on class actions—these claims can be thwarted by arbitration clauses, by choice-of-law issues, and so on, and Article III standing may be particularly likely to present a hurdle.²¹² But in general, liability regimes that allow for this type of *per se* claim to succeed will be conducive to class actions (and class actions will likewise be conducive to the enforcement of these liability regimes).

IV. USING AGGREGATION TO MANAGE PROBLEMS OF GROUP LIABILITY

The previous Part discussed influences on the availability of aggregation to address harms from AI tool use. As it showed, both the kinds of conduct enabled by AI tools and the liability regimes that arise in response will affect if aggregation is a possibility for plaintiffs.

This Part flips the script and examines how the availability of aggregate litigation can influence what options are on the table for the substantive laws

210. See, e.g., Spencer Overton, *Overcoming Racial Harms to Democracy from Artificial Intelligence*, 110 IOWA L. REV. 805, 820–27 (2025).

211. This approach appears to be the path taken by, for instance, Minnesota's statute addressing deepfake intimate imagery. See MINN. STAT. ANN. § 604.32 (West 2025).

212. See *supra* notes 123–133 and accompanying text.

that govern AI tool use. In particular, this Part argues that the existence of private aggregate litigation can make possible certain types of aggregate liability regimes that might otherwise be difficult or impossible to create.

This possibility isn't often discussed because the usual conceit of class actions is that they can only exist where there are already many individual claims.²¹³ In other words, a class is typically seen as emerging from many potentially valid claims put together; the existence of individual claims precedes the class analytically, and individuals' claims remain formally valid without a class action even if a class action is practically necessary to bring them to fruition.

But as AI-related litigation begins to emerge, it is becoming clear that there are likely to be problems in some cases with establishing liability on an individual basis—problems that go beyond the cost/benefit analysis of negative-value claims that we traditionally associate with class actions. This Part diagnoses two types of AI harm that will demonstrate these problems. The first can be called “probabilistic wrongs,” which are wrongs that are statistically discernable at the group level but difficult or impossible to prove at the individual level. The second can be thought of as “aggregate wrongs,” which are wrongs that arise only by combining an actor's conduct with respect to many individuals, where that conduct with respect to any one person would not amount to the same kind of wrong.

Both probabilistic wrongs and aggregate wrongs are likely to arise with the use of AI tools, and some existing case law suggests that they already have. And with both scenarios, aggregate litigation offers important advantages over individual claims. These advantages go beyond the economies of scale that are typically the focus of aggregate litigation, and instead reflect more fundamental questions about proof and the nature of wrongdoing at the group level. As this Part shows, aggregate litigation will allow courts to give at least plausible answers to these questions, providing a key mechanism for implementing liability regimes that otherwise might be difficult to achieve.

A. Probabilistic wrongs

This type of wrong arises where it is possible to demonstrate that some portion of the members of a group were harmed by an AI tool, but it is

213. See, e.g., Marcus, *supra* note 93, at 592–94 (discussing the regulatory and adjectival conceptions of the class action). Even though, as Marcus describes, the “adjectival” conception of the class action emphasizes the individuality of claims more than the regulatory conception, even the regulatory conception of the class action arises more out of the desire to address resource constraints and benefit from economies of scale than the desire to litigate the kinds of probabilistic wrongs and aggregate wrongs described in this section. *Id.*

difficult or impossible to determine which specific members are the ones who were harmed.

Take, for instance, the following scenario, which is based off the case *Newman v. Google*:²¹⁴

Biased Content Moderation: A large tech company runs a platform for user-generated videos. For content moderation, the tech company uses a machine-learning tool that screens recently uploaded videos and automatically determines their suitability. A group of Black content creators sues the company as a class, alleging that the machine-learning tool is biased against videos featuring Black individuals. In discovery, evidence emerges comparing similar videos by Black and non-Black content creators, and finding that Black-created videos are rejected about 20% more often.

The difficulty here can be that, even though it is possible to prove that the algorithm is biased, it may be impracticable or impossible to prove that any particular creator's video was rejected because of that bias. To begin with, assessing each video individually may be cost-prohibitive for all of the reasons that the company uses an algorithm for content moderation in the first place: It takes time, expertise, and judgment to review videos and compare them to a set of rules or standards, particularly in a context like a discrimination lawsuit that may require finding and establishing comparators. The cost of doing this at scale may weigh strongly against certifying a class.

But aside from factors of cost, it may also be the case that there is no real "ground truth" as to which videos were rejected because of the algorithm's bias. Such a causal claim requires some sort of comparative baseline—a demonstration that, in the counterfactual world without the defendant's wrongful conduct, the plaintiff would not have suffered a harm.²¹⁵ To try and establish that baseline, it may be possible to build a non-biased algorithm, or at least a less discriminatory algorithm,²¹⁶ and rerun the videos through that algorithm. But the use of such a system may not be a reliable form of proof. That's because there will often be multiple possible

214. See *Newman v. Google LLC*, 687 F. Supp. 3d 863 (N.D. Cal. 2023). The allegations in the case bear some relation to the hypothetical here in a general sense, but the court dismissed the case for failing after "six opportunities to adequately plead their claims." *Id.* at 865 ("The general idea that YouTube's algorithm could discriminate based on race is certainly plausible. But the allegations in this particular lawsuit do not come close to suggesting that the plaintiffs have experienced such discrimination.").

215. See, e.g., Sergio J. Campos, *The Commonality of Causation*, 46 OHIO N.U. L. REV. 229, 253–58 (2020).

216. See Black et al., *supra* note 141, at 69–71.

algorithms that perform about equally well on the given task across a population of inputs as a whole, but which end up producing different results for any one particular input.²¹⁷ Two different algorithms might both be “neutral” with respect to race at the group level but still treat a specific video differently. Showing that a particular video was rejected by a biased algorithm but not rejected by a neutral algorithm may therefore be a poor form of causal evidence in these circumstances.

Antidiscrimination law is a particularly salient example, but this issue can arise in other contexts as well. The potential for this problem exists wherever it is possible to make relatively confident generalizations about the relationship between inputs and outputs at a population level, but where it is less possible to have confidence about whether any particular individual output is “correct” or biased in some way.²¹⁸ This lack of confidence about particular individual cases could be due either to the possibility that different acceptable models result in different outcomes for the same individual,²¹⁹ or it could be due to some element of process randomization internal to a given model’s operations.²²⁰

Consider, for instance, another scenario:

Self-Dealing Chatbot: A brokerage firm markets products to its customers via an LLM-powered chatbot, which makes recommendations and helps complete transactions. Some of these products earn more profit for the firm, and some earn less profit; the firm is under a legal obligation to act in its customers’ best interests rather than its own. Evidence emerges that, across the firm’s customer base as a whole, the chatbot recommends higher-profit (i.e., self-dealing) products much more than a reasonable baseline established by market surveys and expert witnesses. But even a “profit-neutral” algorithm would recommend the higher-profit products to some customers, and it is difficult to establish which

217. See, e.g., *id.* at 61–67 (discussing model multiplicity); see also Emily Black, Manish Raghavan & Solon Barocas, *Model Multiplicity: Opportunities, Concerns and Solutions*, 2022 ACM CONF. ON FAIRNESS, ACCOUNTABILITY, AND TRANSPARENCY 850 (2022), <https://dl.acm.org/doi/pdf/10.1145/3531146.3533149> [<https://perma.cc/CZ9A-TLTW>] (describing model multiplicity in depth); A. Feder Cooper et al., *Arbitrariness and Social Prediction: The Confounding Role of Variance in Fair Classification*, 38 PROC. AAAI CONF. ON ARTIFICIAL INTELLIGENCE 22004 (2024) (demonstrating this problem with respect to models trained on the COMPAS prison recidivism dataset).

218. See Cooper et al., *supra* note 162, at 4 (noting that it is possible to develop multiple trained models that yield a distribution of possible outcomes for any one individual and creating a problem for labeling any output as the “correct” output).

219. See *id.*

220. See *supra* notes 166–67 and accompanying text (discussing the “temperature” parameter as used in large language models and providing the example of the inconsistent New York City chatbot).

individual customers were steered in the wrong direction. That is because the chatbot's recommendations are path-dependent and have a degree of randomness baked into them: To mimic more realistic human conversation, the chatbot makes individualized responses, and its choice of response is sometimes stochastically altered to imitate the variation in natural human interactions. As a result, even individuals with highly similar circumstances (credit scores, preferences, incomes, etc.) receive different recommendations, and the same individual can be run through the chatbot multiple times with a range of different outcomes.

Here, the difficulty of proving causation at the individual level is distinct from the difficulty in *Biased Content Moderation*. There, the difficulty arose from "model multiplicity"—multiple acceptable models may treat the same individual differently, which may at times make it difficult to establish a neutral baseline for causation purposes.²²¹ That problem could also arise in *Self-Dealing Chatbot*, but there is an additional problem here, which might be called "model inconsistency": the variability within a single model in terms of how it treats individuals. It will be difficult to prove individual harm caused by the model's self-dealing where it is true that (a) even with a lawful, non-self-dealing chatbot, some individuals would be recommended higher-profit products, and (b) on different "passes" through the algorithm, the same individual would sometimes be recommended the higher-profit products and sometimes not. Both model inconsistency and model multiplicity thus raise the same problem: A model can be shown to have a potentially unlawful bias for its outputs overall, but it is difficult to prove that that bias has a causal effect for any particular individual.

The difficulty of proving this type of individual causal connection need not be a problem, depending on the liability regime. It may be that, for instance, merely being subjected to a discriminatory or self-dealing model is a per se basis for liability, of the type discussed above. In *Self-Dealing Chatbot*, that might mean that every customer who engaged with the chatbot would be entitled to some sort of statutory damages, perhaps out of recognition that they did not get the quality of advice that they expected to receive. Or a legal regime could presume causation once bias and harm are both shown—in the *Biased Content Moderation* example, that would permit every Black creator who had a video rejected to have a valid claim, even if the degree of bias shown suggests that bias is only responsible for a portion of the rejected videos.²²² These kinds of possibilities are worth considering

221. See, e.g., Black et al., *supra* note 141, at 62–67 (discussing model multiplicity).

222. This assumes here that a rejected video is adequate to establish harm.

for anyone designing substantive liability regimes in an area where it may be possible to prove that an algorithm was biased, but difficult to prove specific individual causation.²²³

But those approaches have their downsides, too. In particular, they may lead to expansive liability that could strike some as inappropriate in situations where there is bias, but where the number of biased decisions is a small fraction of the number of total decisions. In a biased world, many forms of algorithmic decision-making will reflect some forms of bias.²²⁴ The best response is likely not to reject all algorithmic tools, but rather to assess which types and magnitudes of bias are acceptable given the other goals that these tools may serve.²²⁵ In many liability regimes, we seek to proportion liability to the harm caused by a defendant's conduct. In the world of AI tools, that would suggest being able to assign liability that is proportionate to the number of individuals who are harmed.

Past approaches to aggregate litigation provide some guidance for how to maintain a system of proportionate liability where proving individual causation can be difficult. Aggregate litigation, and mass torts in particular, have long had to deal with what Alexandra Lahav calls "chancy causation."²²⁶ Chancy causation is causation that "can only be attributed probabilistically," making it impossible to prove but-for causation in some kinds of cases.²²⁷ Chancy causation has come up for decades in the context of mass torts.²²⁸ With toxic torts, for instance, it is possible to say that exposure to a chemical increases the risk of developing cancer by a certain percentage, but not to know reliably that a particular individual's cancer was caused by the exposure.²²⁹ As a result, scenarios can arise where it is possible to prove that, say, 20% of a given population's cancer was caused

223. This is in some ways the mirror image of what may be a more familiar type of probabilistic harm from the privacy context, such as the harm that occurs in data breaches. *See, e.g.,* Daniel J. Solove & Danielle Keats Citron, *Risk and Anxiety: A Theory of Data-Breach Harms*, 96 TEX. L. REV. 737, 756–63 (2018). In those cases, a data breach compromises a group of individuals' data, but only some subset of that group may end up experiencing economic loss as a downstream consequence of that data breach, such as through identity theft. *Id.* It may be impossible to identify *ex ante* who will end up suffering such economic loss, as losses can take years to manifest. *Id.* In contrast, here the difficulty is that the injury has already happened, but it is difficult or impossible to identify *ex post* which individual members of the group were injured by the AI tool at issue.

224. Mayson, *supra* note 40, at 2225.

225. *Id.*; *see also* Hellman, *supra* note 7, at 838–40 (discussing the importance of context in determining which types and magnitudes of errors are acceptable).

226. *See* Alexandra D. Lahav, Essay, *Chancy Causation in Tort Law*, 15 J. TORT L. 109, 110 (2022).

227. *Id.* at 109.

228. *Id.*

229. *Id.* at 115–16.

by exposure to a toxic product, but impossible to prove that any particular individual's cancer was caused by that exposure.²³⁰

This parallel suggests one potential path forward: issuing damages based on the overall harm caused by the defendant, then sharing those damages proportionately among the members of the class. In the mass torts context, this proposal was made by David Rosenberg as a way of allowing at least partial recovery for plaintiffs.²³¹ And it could be generalized to cases like *Biased Content Moderation* or *Self-Dealing Chatbot* as well. In *Self-Dealing Chatbot*, for instance, let's say that evidence establishes that the chatbot was about 25% more self-dealing than an acceptable baseline: Only 4,000 people should have been recommended the higher-profit products, but another 1,000 people were steered toward them. The proportionate damages approach would determine what a reasonable amount of damages would be in the aggregate for 1,000 people, and then would give all 5,000 people their proportionate share (20%) of those damages.

This approach has its shortcomings, too, of course. Nobody would be fully compensated; and 4,000 people would be compensated who would not "deserve" it in some sense. But where it is impossible to prove individual causation, this might be the most feasible alternative to not compensating anyone at all, or giving all 5,000 people 100% of the relevant damages via a per se liability theory like the one in *Baker*. It achieves a proportionate level of deterrence by causing the defendant to internalize the costs of all of its wrongful conduct, without over deterring.²³²

Courts have so far been reluctant to take that approach as a matter of judge-made law in the world of mass torts.²³³ But the creation of new liability regimes around AI may provide an opportunity to reconsider it as a method going forward. Proposals for AI regulatory regimes often involve

230. *Id.*

231. See, e.g., David Rosenberg, *The Causal Connection in Mass Exposure Cases: A "Public Law" Vision of the Tort System*, 97 HARV. L. REV. 849, 859 (1984) (discussing a proportional liability regime in which damages are distributed "in proportion to the probability of causation . . . despite the absence of individualized proof of the causal connection"); see also Lahav, *supra* note 226, at 132 (discussing Rosenberg's proposal).

232. Cf. Rosenberg, *supra* note 231, at 866; see also Saul Levmore, *Probabilistic Recoveries, Restitution, and Recurring Wrongs*, 19 J. LEGAL STUD. 691, 697–98 (1990) (discussing the value of probabilistic liability rules in aggregate scenarios).

Another possibility here, which retains the feature of deterrence if not compensation, is the *cy près* award. *Cy près* awards provide a way of distributing damages to people other than those who were injured by a defendant's actions. See Thomas E. Kadri & Ignacio N. Cofone, *Cy Près Settlements in Privacy Class Actions*, in CLASS ACTIONS IN PRIVACY LAW 99 (Ignacio N. Cofone ed., 2021). Such a tool can be valuable in contexts where per-person damages awards might be small, or where it is costly or impossible to determine the damages for specific individuals within a broader group of injured people. See *id.* at 105–07.

233. Lahav, *supra* note 226, at 132 & n.79.

tools such as audits that would be useful for detecting bias at a group level.²³⁴ But as this Section has shown, converting that to individual claims may have some difficulties. Aggregate litigation creates the possibility of empowering a liability regime to operate at the group level while retaining the advantages of private enforcement.

B. Aggregate wrongs

A related, but distinct, form of potential group-based liability can be thought of as “aggregate wrongs.” Here, the problem is not the probabilistic nature of the causation involved, but instead the fact that the causation involved is a kind of aggregate causation: A defendant’s conduct with respect to many individuals adds up to causing a wrong, but the conduct does not amount to a wrong with respect to any particular individual.

Consider the following scenario:

The Minor Authors’ Guild: A large language model is trained by a company on a vast dataset that contains many copyrighted works. The Minor Authors’ Guild, a group of authors who have written only one relatively unsuccessful book each, brings suit as a class action on behalf of themselves and others similarly situated. The authors’ works are contained in the language model’s training data, and their central claim alleges that the company was unjustly enriched from training its model on their copyrighted work.

The plaintiffs succeed in establishing that it was unlawful to train on their copyrighted works, but hit a snag when it comes time to prove the defendant’s enrichment. Unjust enrichment provides a remedy based on the profits attributable to a defendant’s conduct.²³⁵ The defendant’s evidence establishes that, for any given author, the LLM would have been able to function just as well if it had been trained without that author’s work in the dataset, and no commercial revenues would have been different.²³⁶ But the plaintiffs’ evidence

234. See, e.g., Neel Guha et al., *AI Regulation Has Its Own Alignment Problem: The Technical and Institutional Feasibility of Disclosure, Registration, Licensing, and Auditing*, 92 GEO. WASH. L. REV. 1473, 1529 (2024).

235. See, e.g., RESTATEMENT (THIRD) OF RESTITUTION AND UNJUST ENRICHMENT § 51(4) (AM. L. INST. 2011).

236. This evidence, while hypothetical in this scenario, is plausible. State-of-the-art large language models train on vast amounts of data, and unless one author’s work is particularly voluminous or distinctive, there will plausibly be scenarios where no specific revenue (and possibly no specific

establishes that if all of the class members' works had been removed from the training data, the language model would have had a significant drop in quality and would not have been commercially viable.

In this scenario, it would be feasible to conclude that no individual author in this class would have a successful claim for unjust enrichment, because there is no profit attributable to the defendant's conduct with respect to any individual author.²³⁷ But at the same time, the class as a whole accounts for work that is collectively a but-for cause of all of the defendant's profit. This dichotomy is possible because of the nature of the machine learning process: The language model has learned patterns of human language that are present in many works, such that there may be many works that were not individually essential for the model's ability to learn—but collectively, if you add up enough works they amount to a necessary part of the dataset.²³⁸

There are fewer analogues here than in the world of chancy causation, but there are some. The problem here may be described as one of "aggregate causation"—where the actions underlying any individual claim did not cause the defendant's profit, but the actions underlying each claim collectively did cause the profit.²³⁹ Typically, these actions arise in scenarios where there is one plaintiff and multiple people who collectively caused their injury, and the question arises as to who is liable if no defendant's conduct was a necessary or sufficient cause of the injury.²⁴⁰ *The Minor Authors' Guild* presents an inverted scenario, where there is one defendant, and that defendant's wrongdoing becomes a but-for cause of its enrichment only when it is aggregated over many different plaintiffs.²⁴¹ In addition to the intellectual property context, such a problem may arise in the privacy context. For instance, there may be many scenarios where individual persons' private data is not particularly valuable or able to cause harm on their own, but when that data is combined across many persons it becomes

functionality at all) is attributable to the presence of that author's work in the training data. *See generally* Daniel Wilf-Townsend, *The Deletion Remedy*, 103 N.C. L. REV. 1809 (2025).

237. *See, e.g.*, RESTATEMENT (THIRD) OF RESTITUTION AND UNJUST ENRICHMENT § 51(5)(d) (AM. L. INST. 2011) ("A claimant who seeks disgorgement of profit has the burden of producing evidence permitting at least a reasonable approximation of the amount of the wrongful gain.").

238. *See, e.g.*, Wilf-Townsend, *supra* note 236 (discussing how works in a corpus of training data tend to be more valuable to the extent that they are larger and/or more particularly distinctive compared to the rest of the data).

239. *See* Lahav, *supra* note 226, at 118.

240. *See, e.g., id.* (giving the example of multiple chemical exposures, none of which caused cancer acting alone but which cumulatively caused a person's cancer).

241. It is important to distinguish this issue of aggregate causation from the more run-of-the-mill types of aggregation that are typically the justification for aggregate litigation.

a valuable tool for making inferences (and can potentially cause harm as well).²⁴²

At first, the advantages of aggregate litigation in these contexts seem obvious—in that no individual plaintiff has a viable claim, but it seems as if the group as a whole does. But to make aggregate litigation work here, a little massaging is needed for some aspect of the liability regime. Although some doctrines of class actions have aspects of treating the class more as a distinct entity than as a bundle of discrete individuals,²⁴³ there is not generally a doctrine of emergent liability that allows for liability to flow to the class as a whole where no individual member of that group has an individually valid claim.²⁴⁴

To permit claims in this kind of scenario, either a new substantive type of liability may need to be created, or courts may need to be comfortable being creative with their assessment of liability. Where the focus is on a remedy like unjust enrichment, which has a flexible origin in equity, it may be possible to adopt this sort of theory without legislative action. It seems likely that such an adoption would arise only in the context of aggregate litigation, where the claims of many are directly before a court and implicitly or explicitly accepting an aggregative theory is the only way to accord relief.

Where legislation is necessary to create a liability regime involving aggregate causation, aggregate litigation may not be strictly necessary. Other regimes, such as *qui tam* or private attorney general regimes, allow private plaintiffs to vindicate wrongs committed against large numbers of people. These kinds of structures could also conceivably be used to permit private causes of action by individuals under a kind of aggregative theory. But aggregate litigation, and class actions in particular, is already a well-established device for bringing many claims together in one place with a number of safeguards for due process. If aggregate wrongs become enough of a problem that they need their own liability regime to be established,

242. See Solow-Niederman, *supra* note 40, at 382–83 (describing how the value, use, and potential for harm of an individual’s data can change in the context of machine learning tools that have “the ability to discern patterns by analyzing large datasets”).

243. See, e.g., Marcus & Ostrander, *supra* note 137, at 1543 (mentioning cases consistent with an “entity theory” of class actions); Alexandra Lahav, *Fundamental Principles for Class Action Governance*, 37 IND. L. REV. 65, 106–08 (2003) (discussing the entity theory).

244. Class action doctrine allows something closer to this with a class action brought under Rule 23(b)(2), which permits courts to order relief for a group as a whole without inquiring into the details of the claims of each individual member. See, e.g., Maureen Carroll, *Class Actions, Indivisibility, and Rule 23(b)(2)*, 99 B.U. L. REV. 59, 68–71 (2019); Robert G. Bone, *The Puzzling Idea of Adjudicative Representation: Lessons for Aggregate Litigation and Class Actions*, 79 GEO. WASH. L. REV. 577, 610–11 (2011). But even in those cases, the named plaintiffs must have individually valid claims, which would not be true of the hypothetical plaintiffs in the *Minor Authors’ Guild* scenario.

aggregate litigation is the obvious place to look for inspiration as to how, procedurally, such a liability regime could be privately enforced.

CONCLUSION

The world of AI regulation is just the latest source of evidence for the proposition that substance and procedure are never far apart. When it comes to the major goals of public policy—expressing values and shaping conduct—the enforcement of the law must be a primary consideration alongside the law's substance. Nowhere is that more true than in current efforts to govern the use of AI tools. The scale and complexity of AI-related litigation is likely to necessitate aggregate enforcement devices. And the possibilities of aggregate enforcement also provide opportunities for managing different questions about the structure of liability regimes as well. Recognizing the mutual influence between liability standards and enforcement mechanisms will thus help lawmakers both ensure that they are using all of the substantive room available to them, and that their rules for liability are effectively enforced.