



The Journal of Robotics, Artificial Intelligence & Law

Editor's Note: The Future
Victoria Prussen Spears

Flying Cars: Are We Ready for Them?
Elaine D. Solomon

U.S. AI Regulation Guide: Legislative Overview and Practical Considerations
Yoon Chae

Artificial Intelligence in Healthcare: Can Regulation Catch Up with Innovation?
Alaap B. Shah

New York and New Jersey Make an Early Effort to Regulate Artificial Intelligence
Marc S. Martin, Charlyn L. Ho, and Michael A. Sherling

Artificial Intelligence and the Fair Housing Act: Algorithms Under Attack?
William T. Gordon, Katherine Kirkpatrick, and Katherine Mueller

Angel Investing Lessons: The First Mover Disadvantage
Paul A. Jones

The Convertible Debt Valuation Cap: The Trigger Financing Investor Perspective
Paul A. Jones

Big News for Small Mobility: Germany Opens Up to E-Scooters
Andreas Grünwald, Christoph Nüßing, and Theresa Oehm

Everything Is Not *Terminator*: AI Issues Raised by the California Consumer Privacy Act
John Frank Weaver

- 5 Editor's Note: The Future**
Victoria Prussen Spears
- 9 Flying Cars: Are We Ready for Them?**
Elaine D. Solomon
- 17 U.S. AI Regulation Guide: Legislative Overview and Practical Considerations**
Yoon Chae
- 41 Artificial Intelligence in Healthcare: Can Regulation Catch Up with Innovation?**
Alaap B. Shah
- 47 New York and New Jersey Make an Early Effort to Regulate Artificial Intelligence**
Marc S. Martin, Charlyn L. Ho, and Michael A. Sherling
- 53 Artificial Intelligence and the Fair Housing Act: Algorithms Under Attack?**
William T. Gordon, Katherine Kirkpatrick, and Katherine Mueller
- 57 Angel Investing Lessons: The First Mover Disadvantage**
Paul A. Jones
- 63 The Convertible Debt Valuation Cap: The Trigger Financing Investor Perspective**
Paul A. Jones
- 69 Big News for Small Mobility: Germany Opens Up to E-Scooters**
Andreas Grünwald, Christoph Nüßing, and Theresa Oehm
- 73 Everything Is Not *Terminator*: AI Issues Raised by the California Consumer Privacy Act**
John Frank Weaver

EDITOR-IN-CHIEF

Steven A. Meyerowitz

President, Meyerowitz Communications Inc.

EDITOR

Victoria Prussen Spears

Senior Vice President, Meyerowitz Communications Inc.

BOARD OF EDITORS

Miranda Cole

Partner, Covington & Burling LLP

Kathryn DeBord

Partner & Chief Innovation Officer, Bryan Cave LLP

Melody Drummond Hansen

Partner, O'Melveny & Myers LLP

Paul B. Keller

Partner, Norton Rose Fulbright US LLP

Garry G. Mathiason

Shareholder, Littler Mendelson P.C.

Elaine D. Solomon

Partner, Blank Rome LLP

Linda J. Thayer

Partner, Finnegan, Henderson, Farabow, Garrett & Dunner LLP

Mercedes K. Tunstall

Partner, Pillsbury Winthrop Shaw Pittman LLP

Edward J. Walters

Chief Executive Officer, Fastcase Inc.

John Frank Weaver

Attorney, McLane Middleton, Professional Association

THE JOURNAL OF ROBOTICS, ARTIFICIAL INTELLIGENCE & LAW (ISSN 2575-5633 (print) /ISSN 2575-5617 (online) at \$495.00 annually is published six times per year by Full Court Press, a Fastcase, Inc., imprint. Copyright 2020 Fastcase, Inc. No part of this journal may be reproduced in any form—by microfilm, xerography, or otherwise—or incorporated into any information retrieval system without the written permission of the copyright owner. For customer support, please contact Fastcase, Inc., 711 D St. NW, Suite 200, Washington, D.C. 20004, 202.999.4777 (phone), 202.521.3462 (fax), or email customer service at support@fastcase.com.

Publishing Staff

Publisher: Morgan Morrisette Wright

Journal Designer: Sharon D. Ray

Cover Art Design: Juan Bustamante

Cite this publication as:

The Journal of Robotics, Artificial Intelligence & Law (Fastcase)

This publication is sold with the understanding that the publisher is not engaged in rendering legal, accounting, or other professional services. If legal advice or other expert assistance is required, the services of a competent professional should be sought.

Copyright © 2020 Full Court Press, an imprint of Fastcase, Inc.

All Rights Reserved.

A Full Court Press, Fastcase, Inc., Publication

Editorial Office

711 D St. NW, Suite 200, Washington, D.C. 20004

<https://www.fastcase.com/>

POSTMASTER: Send address changes to THE JOURNAL OF ROBOTICS, ARTIFICIAL INTELLIGENCE & LAW, 711 D St. NW, Suite 200, Washington, D.C. 20004.

Articles and Submissions

Direct editorial inquiries and send material for publication to:

Steven A. Meyerowitz, Editor-in-Chief, Meyerowitz Communications Inc.,
26910 Grand Central Parkway, #18R, Floral Park, NY 11005, smeyerowitz@
meyerowitzcommunications.com, 646.539.8300.

Material for publication is welcomed—articles, decisions, or other items of interest to attorneys and law firms, in-house counsel, corporate compliance officers, government agencies and their counsel, senior business executives, scientists, engineers, and anyone interested in the law governing artificial intelligence and robotics. This publication is designed to be accurate and authoritative, but neither the publisher nor the authors are rendering legal, accounting, or other professional services in this publication. If legal or other expert advice is desired, retain the services of an appropriate professional. The articles and columns reflect only the present considerations and views of the authors and do not necessarily reflect those of the firms or organizations with which they are affiliated, any of the former or present clients of the authors or their firms or organizations, or the editors or publisher.

QUESTIONS ABOUT THIS PUBLICATION?

For questions about the Editorial Content appearing in these volumes or reprint permission, please call:

Morgan Morrisette Wright, Publisher, Full Court Press at mwright@fastcase.com
or at 202.999.4878

For questions or Sales and Customer Service:

Customer Service
Available 8am–8pm Eastern Time
866.773.2782 (phone)
support@fastcase.com (email)

Sales
202.999.4777 (phone)
sales@fastcase.com (email)
ISSN 2575-5633 (print)
ISSN 2575-5617 (online)

Artificial Intelligence in Healthcare: Can Regulation Catch Up with Innovation?

Alaap B. Shah*

The healthcare industry continues to grapple with legal and ethical questions about how to responsibly develop, implement, and use artificial intelligence (“AI”). These uncertainties arise in part because there are few laws or regulations that directly address AI technology in general or its application to healthcare specifically. The author of this article explains the issues and risks surrounding the use of AI in healthcare and notes that until regulatory clarity emerges, the healthcare industry will be left to manage risks by creatively applying existing laws and regulations to AI paradigms.

Leveraging artificial intelligence (“AI”) in healthcare is very promising, and has already produced some astounding results in areas such as in radiological imaging. Yet the healthcare industry continues to wrangle with legal and ethical questions about how to responsibly develop, implement, and use such technologies. These uncertainties arise in part because there are few laws or regulations that directly address AI technology in general or its application to healthcare specifically. While the lack of clear boundaries arguably creates a “greenfield” environment for innovation, risks borne by AI in healthcare should be considered for regulation (either *ex ante* or *ex post*). Until regulatory clarity emerges, the healthcare industry will be left to manage these risks by creatively applying existing laws and regulations to AI paradigms.

Regulating AI Remains Complex

Effective regulation of AI remains challenging due to several unique aspects of these technologies and their applications to healthcare. First, AI innovation is affecting a number of diverse segments of healthcare that face different regulatory risks. Innovators are looking to lend AI to areas such as clinical decision support, utilization review, reimbursement and payment, and research, among many others. This makes it difficult to establish

a one-size-fits-all regulatory framework. Second, even within a specific segment of the industry, AI solutions are being created for an array of different purposes, which further frustrates creating a unified approach to regulation. Third, even if the absence of this variability, the definition of AI¹ remains debated among experts and regulators. Without a common understanding about the attributes of these technologies that create risk, stakeholders will have difficulty determining what aspects of AI warrant regulation.² Finally, many AI technologies are developed and operate as “black boxes” with opaque processes, often with the capacity to engage in unforeseeable actions. Without a robust understanding of how these technologies function, regulators will have difficulty developing guardrails for responsible development and use of AI.

Ethical Dimensions for Managing AI Risks

Any effective risk management approach for AI must consider a variety of aspects of such technologies in the context of ethical principles. Chief among these ethical principles are transparency of algorithms and architecture, reliability and fairness of inputs and outputs, accountability features, and safeguards around privacy, security, and safety. Many stakeholders assert that trust in AI solutions may only be achieved if such systems are measurably reliable, transparent, explainable, and able to achieve repeatable results (a combination of features referred to as “interpretability”).³ Further, the interpretability of AI should be designed in a manner to allow for translation of the algorithm’s “reasoning” in to terms that are meaningful for the human end-users (beyond the technologist that created the AI).⁴

Privacy and Data Security Risks Associated with AI

As with many technologies employed in healthcare, throughput of high-quality, high-volume patient data is central to proper functioning. This requirement creates risk to patient privacy, and AI is no exception. AI depends heavily on the collection, transmission, and analysis of large volumes of identifiable patient data that flows through these solutions. It is also important to note that privacy

and security risks arise throughout the AI product life cycle from development to implementation through until end-of-life. Thus, to effectively manage these risks, privacy and security features must be baked in to AI solutions by design from the outset.

Privacy and data security risks first arise when developers seek access to large patient data sets to test and train AI technologies. This requires adequate data rights and compliant data sharing mechanisms to support sharing and use of patient data for development purposes. Even if adequate compliance and data sharing rights exist, data security must also be considered when collecting, storing, and processing of large quantities of patient data.⁵ Once implemented, entities should undertake privacy and security risk assessments and mitigation activities to manage ongoing risks associated with patient data flowing through AI solutions.⁶

Risks may also arise if AI solutions process data in unintended ways. An interesting attribute of AI is that deep learning in neural networks sometimes leads to unanticipated or undesirable behaviors of the AI. As such, it is conceivable that an AI algorithm trained to do one thing with patient data may conduct other activities with such data. Finally, even when an AI technology is decommissioned, patient data stored within the AI solution must still be securely removed.

Safety Risks Associated with AI

Another risk generated by AI in healthcare relates to patient safety. To the extent clinicians rely on AI, such as clinical decision support algorithms, to render medical advice, malpractice risk may be borne. These types of risks may arise due to the quality of data used to train the AI, the quality of the data inputs, the reliability of outputs, and the ability to audit the algorithms rationale. Clearly, AI-supported clinical decision-making will require several layers of iterative vetting to effectively manage risk to patients.

Patient harm may also arise even when a clinician is not making any clinical judgment. For example, implantable or wearable devices that are vulnerable to cybersecurity compromise could impact patient safety. Such compromise through unauthorized access could result in tampering with data integrity or device functionality resulting in patient harm. Likewise, AI solutions that operate independent from human intervention could pose patient safety

risks to the extent these fully automated technologies do not function with adequate levels of accuracy, precision, and accountability.

Bias and Fairness Risks Associated with AI

AI also poses risks related to patient discrimination as well as fraud and abuse arising from unreliability, inaccuracy, or bias associated with outputs.⁷ For example, improper AI training could result in output bias that causes discriminatory coverage or treatment determinations or access to healthcare.⁸ Similarly, improper training or output inaccuracies related to AI-facilitated billing and payment functions could create fraud and abuse compliance risk. Use of AI without properly training algorithms with high-quality data and implementing mechanisms to verify accuracy and precision of outputs over time could pose significant risk associated with bias in outputs.

The Debate About AI Regulation Continues

While a robust AI regulatory approach remains elusive, there is emerging consensus on a few aspects of AI regulation. First, any useful regulatory scheme should foster generation and maintenance of trust in AI technology. Second, any regulation should address foreseeable risks without stifling innovation. Accordingly, regulators should aim to strike the right balance between regulating development and use from a premarket and post-market perspective. Third, given the variability in AI technologies and their applications within healthcare, a sliding scale approach⁹ to regulation is likely the more appropriate approach. This will allow the industry and public to learn about the benefits and how to better manage risks over time.

Closing Thoughts and Next Steps

Innovation in AI is clearly outpacing the law. Unlike many other areas of law, where clear legal and regulatory schemes exist, AI technologies are being developed and adopted without many concrete guardrails. Yet, as with any new, disruptive technologies, successful adoption of such technologies depends on trust.

Unfortunately, humans generally have difficulty establishing trust with things we cannot sufficiently control, effectively manage risk around, or adequately comprehend. As such, humans continue to have “trust issues” given the nature of AI. So what are we to do?

In the immediate term, navigating these issues without direct legal or regulatory schemes will require addressing risk by evaluating AI tools based on policy and ethical principles such as reliability, safety, transparency, accountability, and fairness. In the intermediate term, early adopters of AI will be left to allocate risk through insurance and contracts with developers, partners, and third parties leveraging AI technology. In the longer term, expect to see government regulators increasingly make efforts to evaluate AI risk and generate new regulatory approaches to manage AI risks.

Notes

* Alaap B. Shah is a member of the firm Epstein Becker & Green, P.C., in the Health Care and Life Sciences practice, focusing on defense and counseling of health care entities on legal and regulatory compliance issues around privacy, cybersecurity, and data asset management. He may be reached at abshah@ebglaw.com.

1. See, e.g., S.J. RUSSELL & P. NORVIG. *ARTIFICIAL INTELLIGENCE: A MODERN APPROACH* at 2 (3d ed. 2010).

2. See P. Stone et al., *Artificial Intelligence and Life in 2030. One Hundred Year Study on Artificial Intelligence: Report of the 2015-2016 Study Panel*, Stanford University, Stanford, CA, at 48 (September 2016) available at <http://ai100.stanford.edu/2016-report> (last accessed Jan. 30, 2019).

3. See T. Zarsky. *Transparent predictions*. 4 UNIV. ILL. L. REV. 1503, 1520 (arguing “interpretability could be considered as an important step to assure the prediction process’s quality, precision, and that the results it provides are not merely anecdotal”).

4. C. Reed. *How should we regulate artificial intelligence?* 376 PHIL. TRANS. R. SOC. A at 7 (2018).

5. See, e.g., SANS Institute. *Making Database Security an IT Security Priority* (Nov. 2009) at 11-12 available at <https://software-security.sans.org/resources/paper/reading-room/making-database-security-security-priority> (last accessed Jan. 30, 2019).

6. See generally, A. Shah, *Death by a Thousand Cuts: Cybersecurity Risk in the Health Care Internet of Things*, AHLA Weekly (May 19, 2018).

7. See, e.g., A.I. Could Worsen Health Disparities, N.Y. Times (Jan. 31, 2019) available at <https://www.nytimes.com/2019/01/31/opinion/ai-bias-healthcare.html> (last accessed Jan. 31, 2019).

8. *Id.*; see also R. Caruana et al., *Intelligible models for healthcare: predicting pneumonia risk and hospital 30-day readmission*. In Proc. 21th ACM SIGKDD Int. Conf. on Knowledge Discovery and Data Mining, Sydney, Australia at 10–13 (Aug. 2015), pp. 1721–1730 *available at* <http://people.dbmi.columbia.edu/noemie/papers/15kdd.pdf>.

9. Scholars have posited various sliding scale approaches to regulation of these technologies. See, e.g., *How should we regulate artificial intelligence*, 376 PHIL. TRANS. R. SOC. A at 8–10 (2018) (suggesting that the law should demand varying degrees of transparency depending on the level of benefit to society, level of harm to individuals, and whether such harm would be legally compensable); A. Rao and E. Cameron, *The Future of Artificial Intelligence Depends on Trust*, STRATEGY & BUS. (July 31, 2018) *available at* <https://www.strategy-business.com/article/The-Future-of-Artificial-Intelligence-Depends-on-Trust?gko=af118> (last accessed Jan. 31, 2019) (suggesting that a sliding scale regulatory approach requiring transparency should depend on the level of risk posed and the level of involvement of a human operator).