

Convener on Artificial Intelligence and Health Law

Designing a Trusted Framework for the Application of AI in Health Care



Convener on Artificial Intelligence and Health Law: Designing a Trusted Framework for the Application of AI in Health Care

Carol Eoannou, Former Editorial Director, Health and Litigation, Bloomberg Law

Convener workgroup contributors:

Bernadette Broccolo, McDermott Will & Emery LLP

Leeann Habte, Best Best & Krieger LLP

Kristen Rosati, Coppersmith Brockelman PLC

Alaap Shah, Epstein Becker & Green PC

This paper reflects the views, recommendations, and conclusions of the Authors listed above and was informed by a discussion with the Convener Participants listed on the following page. The views of the Authors are their own and do not reflect a consensus of opinions, perspectives, conclusions, and recommendations of the Convener Participants. The opinions and perspectives of the Convener Participants are their personal views and do not reflect those of the organizations with whom they are affiliated with or employed by.

Convener Participants

Elise Anthony, Executive Director of Policy, HHS Office of the National Coordinator for Health Information Technology

Daniel Barth-Jones, Assistant Professor of Clinical Epidemiology, Columbia University Mailman School of Public Health

Elisabeth Belmont, Corporate Counsel, MaineHealth

Stephanie Cason, Product Counsel, Google Health

Glenn Cohen, James A. Attwood and Leslie Williams Professor of Law, Harvard Law School, Faculty Director, Petrie-Flom Center for Health Law Policy, Biotechnology & Bioethics

Nathan Leong, Senior Attorney, Director US Health & Life Sciences, Microsoft

Michael Lipinski, Division Director of Federal Policy & Regulatory Affairs, HHS Office of the National Coordinator for Health Information Technology

Belinda Luu, Strategic Leader of Enterprise Data Management and Governance; Senior Counsel at Kaiser Foundation Health Plan

Ryan Mehm, Federal Trade Commission, Division of Privacy and Identity Protection

Bakul Patel, Director, Digital Health Center of Excellence, FDA

Nicholson Price, Professor of Law, University of Michigan School of Law

Jack Resneck Jr., MD, Immediate Past Chair, Board of Trustees, American Medical Association

Ingrid Vasiliu-Feltes, MD, Chief Quality and Innovation Officer, MEDNAX

Convener Planning Workgroup/Moderators

Bernadette Broccolo, McDermott Will & Emery LLP

Leeann Habte, Best Best & Krieger LLP

Jackie Olson, Apple

Kristen Rosati, Coppersmith Brockelman PLC

Alaap Shah, Epstein Becker & Green PC

Introduction:

The American Health Law Association (AHLA) hosted a virtual Convener¹ on Artificial Intelligence (AI) and Health Law on November 2, 2020. AHLA Conveners gather thought leaders including regulators, clinicians, private practitioners, and other leading authorities to problem-solve through candid dialogue and frank discussion on emerging issues that impact health care and health law. Although AHLA does not engage in lobbying activities, topics discussed during AHLA Conveners frequently present a path forward to the discussion of thorny problems confronting the health care industry. In light of AI's novel technical nature, its myriad potential applications, and its dependence on complex big data strategies, it is a particularly appropriate topic for the in-depth focus of an AHLA Convener.

The Convener participants and planning committee members have extensive expertise in big data and health care. They were tasked with identifying the elements of a trusted framework for implementing AI in the health care industry. Their positions in health systems, government, academia, private business, legal practice, clinical medicine, statistics and health information, and consumer technology provided a broad spectrum of multi-disciplinary expertise and perspectives needed for a robust and well-informed discussion. The opinions and perspectives expressed by the participants, are theirs alone and do not reflect those of the organizations with whom they are affiliated or employed by.

This white paper summarizes the Convener discussion, which focused primarily on data privacy and security, regulation, liability allocation, intellectual property, and contracting challenges, and how these issues affect the development of a trusted framework for using AI in health care. It also provides one Author's view of significant regulatory actions taken in the interim between the Convener and the paper's publication.

Notable Developments Since the Convener

Two important federal publications relevant to this discussion were released subsequent to the Convener. First, the Department of Health and Human Services (HHS) published an AI strategy for HHS's approach to using and regulating AI.² HHS established an AI Council to lead HHS's AI's strategy and key strategic priorities.

Second, on January 12, 2021, the Food and Drug Administration (FDA) released a plan for its approach to regulating medical software that contains AI and machine learning components.³ The FDA intends to:

- ▶ Develop an update to the proposed regulatory framework presented in the AI/ML-based SaMD [Software as a Medical Device] discussion paper,⁴ including through the issuance of a Draft Guidance on the Predetermined Change Control Plan;
- ▶ Strengthen FDA's encouragement of the harmonized development of Good Machine Learning Practice (GMLP) through additional FDA participation in collaborative communities and consensus standards development efforts;
- ▶ Support a patient-centered approach by continuing to host discussions on the role of transparency to users of AI/ML-based devices. Building upon the October 2020 Patient Engagement Advisory Committee (PEAC) Meeting focused on patient trust in AI/ML technologies, hold a public workshop on medical device labeling to support transparency to users of AI/ML-based devices;

1 This Convener was initially scheduled as a live event in Washington, D.C. on March 10, 2020. The COVID-19 pandemic necessitated its delay and ultimate presentation via Zoom.

2 <https://www.hhs.gov/sites/default/files/final-hhs-ai-strategy.pdf>.

3 <https://www.fda.gov/media/145022/download>.

4 <https://www.fda.gov/media/122535/download>.

- ▶ Support regulatory science efforts on the development of methodology for the evaluation and improvement of machine learning algorithms, including for the identification and elimination of bias and the robustness and resilience of AI algorithms to withstand changing clinical inputs and conditions; and
- ▶ Advance real-world performance pilots in coordination with stakeholders and other FDA programs, to provide additional clarity on what a real-world evidence generation program could look like for AI/ML-based SaMD.

Clearly, this field is moving quickly. AHLA is providing many educational opportunities to help its members and the public stay abreast of the developments in AI in health care.

AI Opportunities and Challenges

AI is emerging as an area of intense interest in health care. It is expected to have wide-ranging, significant impact in clinical care, research, and health care operations, including the potential to:

- ▶ improve early diagnoses and prevention of a variety of conditions;
- ▶ enhance imaging accuracy and efficiency;
- ▶ fuel the further development and implementation of precision medicine;
- ▶ expand telehealth and virtual care;
- ▶ streamline and expand access to clinical trials;
- ▶ predict demands on emergency department resources and in-patient volumes;
- ▶ fuel patient and consumer engagement and satisfaction;
- ▶ enhance population health management capabilities;
- ▶ reduce provider burn-out;
- ▶ manage supply chains; and
- ▶ improve revenue cycle management opportunities.

Realizing the full potential of these and other applications of AI in the health care context will require the trust and confidence of patients, providers, consumers, researchers and the general public. Earning and maintaining such trust will require a multi-disciplinary framework for development and implementation that supports the “trustworthiness” of the AI solutions.

Building such a framework will not be easy. It will require the time, effort and thought leadership of the many disciplines represented by the convener participants to overcome significant challenges such as:

- (a) the lack of a universal definition and understanding of what constitutes AI in the context of health care and life sciences across the broad and complex spectrum of AI technology functionality sophistication;
- (b) an existing legal and regulatory scheme that is difficult to apply in this context because it does not contemplate the use of AI in health care and is undergoing change at a pace that lags significantly behind the rapid pace of the development and demand for AI;
- (c) protecting the privacy and integrity of data in the vast, complex and ever-evolving personal data ecosystem that is essential for the development and deployment of accurate and unbiased AI algorithms;
- (d) establishing and allocating accountability, risk and liability among the myriad stakeholders and collaborators involved the lifecycle of an AI solution who have varying goals and objectives, compliance obligations and enterprise risk posture;

- (e) identifying, protecting and allocating intellectual property rights under current and developing intellectual property laws; and
- (f) identifying and managing the risk of social, racial and other bias and inequity in AI development and implementation.⁵

Key Characteristics and Elements of a Trusted Framework for the Development and Deployment of AI

The Convener participants discussed potential key characteristics and elements of a trusted framework for the development and deployment of AI in health care. They discussed that a trusted framework:

- a. Encompasses all of the key areas of legal and regulatory risks covered in the Convener;
- b. Strikes an appropriate balance between and among managing legal and regulatory compliance risks, maintaining flexibility to adapt to the dynamic nature of AI innovation, and fueling innovation;
- c. Establishes shared accountability and responsibility among all stakeholders involved throughout the entire development and deployment life cycle;
 - i Shared accountability must be built into both the legal and regulatory framework and contractual relationships among the collaborators.
 - ii Regulatory framework should:
 - 1. Apply to the full spectrum of stakeholders who control risk and derive value from the AI initiative,
 - 2. Align the degree of risk the AI solution presents with an acceptable degree of the risks targeted by applicable regulatory compliance requirements,
 - 3. Be harmonized between and among federal and state laws to the greatest extent possible (consider pre-emption of federal over state) that addresses, among other things:
 - a. De-identification⁶ standards, and
 - b. Prohibition re-identification of de-identified data by any stakeholders.
 - 4. Require a scheme for ongoing assessment of the AI solution's compliance with legal and regulatory requirements (e.g., ongoing testing and re-validation of safety, efficacy and privacy and security protections),
 - 5. Require a robust data governance/stewardship regime and provide guidelines and standards for development of same,
 - 6. Encompass an enhanced combination of technological and administrative requirements and standards for addressing the risks addressed by the legal and regulatory scheme, and
 - 7. Address novel Intellectual Property dimensions presented by AI.
- d. Align contractual allocation of risk with the applicable stakeholders' ability to control/mitigate the risk and the value they derive from the development and deployment of the AI solution.

⁵ While equity considerations were recognized generally as an essential consideration, time constraints prevented full, independent examination of such considerations during the Convener; they were, rather, acknowledged as generally informing all of the topics explored.

⁶ Note: the terms "de-identified" and "de-identification" are used throughout this discussion. That term derives from the Health Insurance Portability and Accountability Act (HIPAA). Many of the issues, observations and recommendations in this discussion of de-identification are also applicable to the concepts of "anonymization" under the EU's General Data Protection Regulation (GDPR) and other domestic and international privacy laws.

Privacy and AI: Recognizing the Need for a Robust Regime for De-Identification of Data and Prevention of Re-Identification

Many AI algorithms perform functions such as clinical care decision support, revenue cycle management, analytics to support monitoring and improving the quality of care and patient outcomes, and streamlining the research study data recruitment and operations. AI algorithms also use identifiable data to create de-identified data and to provide monitoring and detection of re-identification attempts.

Regardless of the specific functionality of the AI technology, AI is fundamentally dependent on big data for both its development and deployment. The dependency makes data privacy protection an essential component of any framework of trust for developing and using AI. The feasibility of data de-identification⁷—the process used to prevent an individual’s personal identity or other identifiable personal information from being revealed—is central to this privacy component, as is the risk of re-identification—the practice of matching de-identified or anonymous data with other available information in order to discover the individual to which the data belong. The ability to de-identify data and to prevent re-identification will be particularly pertinent to the use of identifiable data for initial development of the AI technology and for ongoing training and enhancement of the technology.

In any context, the life span of a de-identified data set or de-identification methodology, and the associated re-identification risk, can vary significantly based on various factors, such as:

- ▶ the quantity of the underlying data;
- ▶ the sources of the data (e.g., several different data sets, consumer wearables and mobile devices and apps);
- ▶ the form of the data (e.g., structured or unstructured [e.g., free text notes in electronic medical records, radiological images, photographs or videos, consumer-generated data]);
- ▶ the type of data (e.g., biometric data, genomic data, other sensitive data); and
- ▶ whether the de-identified data will over time be combined or used in conjunction with other identifiable or de-identified data sets.

The use or creation of de-identified data by an AI algorithm itself has the potential to increase that variability.

Such variability and associated risk will require ongoing monitoring, re-assessment and management of the de-identification process and associated re-identification risk over the life cycle of the technology that is designed to account for the specific current and future uses of the AI technology, changes in the technology itself over time, current and potential future sources and types of data, and the de-identification technique used. Obtaining opinions that statistically certify to data de-identification is a recognized privacy compliance pathway under HIPAA, and it will likely be an increasingly important, if not essential, element of the privacy dimension of the framework of trust for development and deployment of AI solutions, both for purposes of the initial de-identification and ongoing re-validation. With regard to AI solutions that perform data de-identification, an important question is whether the de-identification opinion can certify the algorithm itself, instead of or in addition to certifying a statistically significant sampling of data sets produced by the algorithm.

⁷ For entities covered by HIPAA, Standards for Privacy of Individually Identifiable Health Information, 45 C.F.R. pts. 160 and pt. 164, subpts. A and E, deidentifying patient data can be accomplished either by (a) removing the specific identifiers listed in that standard (which include both personal identifiers and demographic data), or (b) obtaining an opinion from a qualified expert that that the information could be used, alone or in combination with other reasonably available information, by an anticipated recipient to identify an individual who is a subject of the information. See, “More Data Please! The Challenges of Applying Health Information Privacy Laws to the Development of Artificial Intelligence,” *AHLA Connections*, February 2020, available at <https://www.americanhealthlaw.org/publications/health-law-hub-current-topics/artificial-intelligence-and-health-law>.

There is no universal method for measuring re-identification risk. Generally speaking, however, technology tools alone will not be sufficient to achieve the “very small risk” of re-identification threshold contemplated by HIPAA. Therefore, de-identification opinions typically take into account both technical and non-technical risk management measures to assess the re-identification risk. Examples of non-technical measures include:

- ▮ clear identification of anticipated recipients of the de-identified data and limitations or prohibitions on adding recipients;
- ▮ a contractual commitment by the original data recipient and all subsequent recipients not to attempt to re-identify the data;
- ▮ internal administrative bans against intentional re-identification efforts; and
- ▮ mechanisms for enforcing compliance with the conditions upon which the expert determination is based by all data recipients.

The creation and use of “synthetic data” has emerged in recent years as a proxy for real data that provides massive amounts of data for use in training artificial intelligence and machine learning models. In concept, synthetic data has the same statistical properties as known datasets but “rows of observations in synthetic datasets do not correspond to identifiable individuals (rows of data) from the original dataset;”⁸ in other words, there is no “1:1 mapping to real data.”⁹ As such, synthetic data by its nature is believed by some either to resolve the privacy risk or to provide enhanced privacy protection.

Nonetheless, whether synthetic datasets will be susceptible to re-identification remains an unsettled issue and the subject of discussion by leading authorities and commentators.¹⁰ For example, as one Convener participant observed, even when no specific individual data are used, using the same *statistical distributions* as real data may result in data about synthetic individuals that may be so similar to data about real world individuals that re-identification from synthetic data may be possible.

Protective approaches to de-identification and associated re-identification risk may also differ depending on the method or model used and whether synthetic data is generated from actual/real datasets, generated fully using statistical models or processes (e.g., simulations) and without use of actual/real datasets, or generated using a hybrid of those two methods.

Participants discussed potential changes to the current legal and regulatory framework to enhance non-technical measures for managing the re-identification risk. For example, the participants discussed the need for a statutory re-identification prohibition applicable to the broad spectrum of stakeholders involved in the chain of trust.¹¹ Whether the legislative change should be enacted at the state or federal level is subject to debate. Leaving it to the states, however, will perpetuate the current lack of harmonization between and among state health care privacy and confidentiality laws and between state and federal laws.

Another important data-related consideration is the quality, accuracy and utility of the data for the intended purpose, which may be affected in whole or in part by the quality, accuracy and utility of any datasets used to create the data. Key factors affecting quality, accuracy and utility include:

- ▮ the potential for bias or discrimination in the underlying data;
- ▮ the exclusion or censoring of certain outliers in the underlying data;

8 Foraker et al., “Are Synthetic Data Derivatives the Future of Translational Medicine”, JACC: Basic to Translational Science, Vol. 3, No. 5 (October 2018).

9 Emam, “Accelerating AI with Synthetic Data”, available at <https://iapp.org/news/a/accelerating-ai-with-synthetic-data/>.

10 Near and Darais, “Differentially Private Synthetic Data”, Cybersecurity Insights (a NIST blog) (May 3, 2021), available at <https://www.nist.gov/blogs/cybersecurity-insights/differentially-private-synthetic-data>, which points to the additional privacy protection offered by applying differential privacy to the synthetic dataset; Emam, *supra* at note 10.

11 See, e.g., Title 1.81.5. California Consumer Privacy Act of 2018 §1798.140 (h)(1)-(4).

- ▮ the ability to generate fully synthetic data to address bias, discrimination or missing data points needed reliably to use the synthetic data for the intended purposes;
- ▮ whether the same data source will be sufficient for all intended uses; and
- ▮ whether and how frequently the quality, accuracy and utility of the database should be reassessed and revalidated, and what methods will be used to undertake such revalidation (including, without limitation, AI technology itself).

Effect of the Information Blocking and Interoperability Rules on Availability of Data for Development and Use of AI Technology

The final Information Blocking Rule¹² issued by the Office of the National Coordinator (ONC) for Health Information Technology pursuant to the 21st Century Cures Act¹³ (the Cures Act) contains provisions pertinent to the availability of the data upon which AI greatly depends. The most notable for purposes of the Convener discussion are the prohibition against information blocking that is likely to interfere with the access, exchange, or use of health care information¹⁴ and the electronic health record certification requirements for health IT developers.

Although focus tends to fall most heavily on information blocking, there are corollaries between the blocking provision and the certification process requirements. The blocking prohibition applies to certain health care providers, developers of certified health IT, or health information networks or health information exchanges.¹⁵ The interoperability requirements are designed to promote the secure exchange of electronic health information with, and use of electronic health information from, other health information technology without special effort on the part of the user.¹⁶ While a discussion of these Cures Act provisions was beyond the scope of the Convener, they were recognized as recent legislative and policy changes that will further the quest for the development and enhancement of the data ecosystem that is needed to fuel AI innovation.

Notably, the Final Rule also outlines certain best practices regarding privacy policies, which include obtaining consent to sell or share certain data, and writing privacy policies in plain language in a manner intended to be informative. It also addresses the ability of providers to educate patients about their choice to share data with mobile apps, but makes clear that such educational efforts do not constitute blocking. Further, the Final Rule does not *require* providers to undertake these educational roles. It does specify, however, that when providers opt to do so, the education must be factually accurate, objective, unbiased, related to the privacy or security issue at hand, and nondiscriminatory.

It has been suggested that patients will choose to release their data to companies that adequately protect their privacy. Relying on such market forces to protect consumer data in the absence of a formal, well-developed legal and regulatory scheme, however, has not been universally accepted as effective in any sector, and may be particularly problematic from a privacy protection perspective in health care, given the sensitive nature of the data at issue, the inability for patients to readily assess the security controls and safeguards of a mobile app (especially since mobile app vendors generally are not governed by HIPAA), and the time it will take for the market to fashion compliant solutions.

12 85 Fed. Reg. 7,0064 (effective on December 4, 2020 except for 45 CFR 170.401, 170.402(a)(1), and the amendments to 45 CFR part 171 which are effective on November 4, 2020) (to be codified at 45 C.F.R. pt. 170 and 171)

13 P.L. 144-255, Approved December 13, 2016, 130 Stat. 233.

14 45 C.F.R. § 171.03(a)(1). Section 4004 of the Cures Act defines practices that constitute information blocking and authorizes the Secretary of Health and Human Services to identify exceptions, or reasonable and necessary activities that do not constitute information blocking.

15 45 C.F.R. § 171.102.

16 Cures Act, §4003.

Cybersecurity: Protecting Large and Complex Data Ecosystems Supporting the Development and Deployment of Health Care and Life Sciences AI Solutions

Data security is a familiar and continuing concern in a big data environment. Factors that affect the feasibility of successfully protecting data in the context of AI include the types, form, source, or quantity of the data, the hardware and software systems within which data is maintained, and processing, as well as the type or complexity of the AI algorithm employed in such processing.

Responsible sharing of large quantities of data and its use in the training of AI is critical to the success of AI. But even if all privacy concerns are resolved, if AI systems and the processes by which they are designed, developed, and deployed aren't adequately protected and if oversight controls are not implemented, AI inputs and outputs may be negatively impacted through compromise of the integrity, availability and confidentiality of data the systems. Such threats may result in data breach, data loss, bias, patient safety risk, negative patient outcomes, inaccurate or unreliable data analytical outputs, privacy risk related to re-identification, and ultimately erosion of consumer trust.

When evaluating security risks impacting health care related AI, several key questions emerge. Does the increased use of AI create new and/or different cybersecurity risks, given the quantity of the data implicated and its sensitivity? Does the very nature of AI present novel security risks that require novel protections? These considerations may impact the proper functioning of AI at various points of its lifecycle from the IT infrastructure design and deployment phases, data input and training phases, and through the data processing and output phases.

AI can be considered a software algorithm that guides functioning of other IT systems components, whether as a medical device or other product. From that perspective, generally recognized cybersecurity concerns germane to IT systems are not unique to AI, such as the need for access and authorization controls, auditing and perimeter security.

However, with regard to security aspects of the AI software algorithms themselves, different sources for risk do exist. Specifically, proper functioning of AI is predicated on testing and training the AI prior to go-live. Certain cybersecurity attacks during testing phases (referred to as "evasion attacks") can attempt to lead the AI to render incorrect outputs.

Likewise, certain attacks target the training phase due to the need to train AI on large and diverse data sets. This creates a novel risk whereby a cyberattack aimed at compromising the integrity of the training data (referred to as a "poisoning attack") can undermine the functioning of the AI and its outputs. For example, introduction of spoofed training data or introduction of a virus that modifies data can ultimately cause "mis-education" of the algorithm that may result in unintended functioning and outputs. Therefore, ascertaining whether the training data used is real and from a trusted source is a useful step at the outset of developing AI.

For AI algorithms that are designed to engage in continuous learning as the AI is used, a concern was expressed that bias may be introduced into AI functioning through exposure to fake or otherwise manipulated data and can influence and vary how the tool learns and operates. Manipulating the training cycles of AI can thus result in downstream risks associated with reliance on the outputs of the AI, including anything from patient safety risk if AI outputs are relied on for clinical decision making, to regulatory risk if AI is relied on when submitting health care claims, to business risks if AI outputs are relied on to make financial decisions.

To preserve data integrity in AI training, controls should be implemented to reduce cybersecurity risks that could result in tainting a data set to produce bias. These controls could include data integrity checks auditing AI training and other routine quality assurance mechanisms, perhaps even involving human oversight where the magnitude of risk is severe (e.g., patient safety).

Cybersecurity Regulatory Standards:

When a company is collecting data for use with AI, it is crucial to have data security protections in place given the large volume of sensitive health data required to train and use AI. The predominant set of security requirements impacting many health care entities that may develop or use AI exist under HIPAA. These include, *inter alia*, controls related to access, authorization, data integrity, encryption in transit and at rest, as well as robust auditing.

In instances, where an entity is not subject to HIPAA, consumer protection laws enforced by the Federal Trade Commission (FTC) may govern.

The FTC's general data security guidance derived from enforcement actions is tech neutral and therefore applicable to AI. Enforcement decisions provide a good jumping-off point for companies needing guidance on how to protect data. A review of FTC's enforcement actions and position statements reveals several some basic security principles favored by the FTC:

- ▶ Don't over collect. If data isn't collected in the first place, obligations to safeguard it are avoided. Focus on getting only the data needed for purposes of achieving the immediate desired goal.
- ▶ Bake privacy and security protections in at the outset of AI product development, verify they work, and test them throughout the process for vulnerabilities.
- ▶ Once the product or application is up and running, keep appropriate restrictions, like firewalls, technology controls and administrative controls in place around the data.

The FDA has also provided guidance on AI related security. The FDA's authority relates to evaluating the safety and efficacy of medical devices (including AI that qualified as Software as a Medical Device or "SaMD"). Accordingly, the FDA has recognized the importance of managing AI security risks by issuing two sets of guidance. First, the FDA issued voluntary post-market guidance addressing cybersecurity for connected medical devices. The FDA also issued guidance on premarket submission related to medical device cybersecurity management. The FDA has also collaborated with the MITRE Corporation to issue a cybersecurity playbook to guide providers in their risk management of medical devices. Collectively, these publications aim to establish trust around the transparency, accountability and reliability of medical devices, such as AI and other SaMD technologies.

In addition to the current regulatory landscape governing development and use of AI, some pending federal legislation has suggested that the National Institute for Standards and Technology (NIST) should become the governing body for developing security standards for AI. In the past, NIST has issued sound guidance around security interests and functioned as an effective partner for both the FTC and FDA. As a technical standards body, NIST could be a good complementary source for basics and setting out a path for getting to common technology practices regarding AI security.

Transparency as a Key Element of Data Stewardship

Companies that develop AI tools or use AI in their business, whether or not the business is subject to HIPAA, should consider the data stewardship and protection principles by which they operate. A key data stewardship consideration is transparency with patients and consumers about the collection and use of large and sometimes highly sensitive data in connection with the development and ongoing deployment of AI. Without transparency, it is difficult if not impossible to create a trusted process. Companies should evaluate how best to be open and transparent and be careful not to deceive consumers about how the information is being collected, for what purposes, how it will used, and with whom it may be shared.

Transparency serves an important compliance risk management function. In particular, Section 5 of the Federal Trade Commission Act (FTC Act)¹⁷ prohibits unfair or deceptive practices in or affecting commerce. Specifically, it provides:

- ▶ An act or practice is unfair if it causes or is likely to cause substantial injury to consumers that is not reasonably avoidable by consumers themselves and is not outweighed by countervailing benefits to consumers or competition.¹⁸
- ▶ A deceptive practice is a material misrepresentation, omission or practice that is likely to mislead a consumer acting reasonably in the circumstances.¹⁹

At the most basic level, deceptions are lies, and if a company lies about its use of patient data, it could find itself within the crosshairs of the FTC.

For a host of regulatory compliance and liability reasons, it is essential that, given their complexity, AI algorithms must be based on a sound, robust methodology that is explainable to the consumer or patient to the extent possible. If the patient or consumer is denied something of value, including health care, based on an algorithmic decision, the regulated entity must be able and willing to explain the reason for the denial. These explanations may be legally required under the Fair Credit Reporting Act²⁰ (FCRA), which protects information collected by various entities, including medical information companies. FCRA's broad scope regulates decisions of such entities that impact consumers' ability to obtain health and life insurance, among other services. AI can be used to create models that feed information on which the covered companies base their decisions. Furthermore, under the FCRA, they are required to notify consumers when an adverse action is taken based on their reporting. This requirement presumably extends to the choice of algorithm on which the reporting decision rests.

A similar analysis applies to the operation and the FTC's enforcement of the Equal Credit Opportunity Act.²¹

Regulating the Development and Use of AI: Establishing Shared Accountability and Responsibility Among Stakeholders Throughout the Development and Deployment Life Cycle

The FDA and FTC are the federal agencies most active in shaping a regulatory scheme for AI in health care.

The FDA:

The FDA's regulatory authority over AI that qualifies as SaMD can be grouped into two key categories. The first category relates to fixed or "static" AI that applies a rules-based logic that does not change over time. By contrast, the second category involves evolving machine learning processes that evolve over time such that the algorithm is not fixed or static.

Knowledge about the difference between the rules based and machine-learning approaches to AI is helpful in understanding the FDA's regulatory construct.

A rule base consists of a large set of rules, composed by one or more subject-matter and rules-construction experts, to separate documents meeting the criteria for relevance, from those that do not. These rules may take the form of complex Boolean queries—specifying the order of words and their proximity—that indicate relevance or nonrelevance, exceptions to those rules, exceptions to the exceptions, and so on, until all the documents in the collection are correctly classified. Supervised machine learning, on the other hand, infers,

¹⁷ 45 U.S.C. § 45(a)(1).

¹⁸ 45 U.S.C. § 45(n).

¹⁹ F.T.C. Policy on Deception, appended to *Cliffdale Associates, Inc.*, 103 F.T.C. 110, 174 (1984).

²⁰ 15 U.S.C. §§ 1681-1681x.

²¹ 15 U.S.C. §§ 1691-1691f.

from exemplar documents, characteristics that indicate relevance or nonrelevance, and uses the presence or absence of those features to predict the relevance or nonrelevance of other documents.²²

The FDA's enforcement framework recognizes that AI occupies a rapidly evolving product space, covering the spectrum from the point where rules-based engines drive outcomes and create solutions, to the midstream, where empirically derived algorithms use data to create solutions. At both these junctures, human intervention continues to be important for validation. The far end of the product evolution spectrum implicates continuous machine learning with minimal human intervention, but questions remain whether and when human oversight may be required for quality assurance purposes across the spectrum of products.

It should be noted that the spectrum of rules-based to machine learning aligns with the spectrum of risk approach that determines the degree of and approach to regulation by the FDA.

Enforcement under this framework focuses on the full life cycle of the product by evaluating who is doing the developing, how the product is being made, and how to address changes that will be made to the product and the process over time. Accordingly, building process and product confidence is going to require continuous rather than episodic FDA oversight. FDA recognized that building trust is not just about getting a product to market. Rather, trust will come from FDA's review of an AI company itself along with its processes and products in order to gain a degree of transparency about product design, results and performance. This will also help FDA and AI companies to communicate about developments over time with end users as the AI products and processes evolve.

Change plans will be helpful in this regard. They will function much like traditional business plans, and (1) recognize that businesses and products will evolve depending on stage of the product's life cycle and (2) document anticipated reactions to those changes.²³

For purposes of a trusted framework for AI in health care, ongoing testing and re-validation of safety, efficacy and privacy and security protections will be part of the scheme required for assessment of the AI solution's compliance with the FDA's legal and regulatory requirements.

Although performance aspects have not been completely resolved with regard to metrics like the number of data points necessary for training algorithms, the questions have at least been identified, particularly at the pre-market threshold, even if the ultimate answers remain unclear.

Transparency is also necessary in managing people, *i.e.*, telling users, whether they are end users or clinicians, what to expect from product performance over time. Preplanning both the nature of the changes and how they will be communicated to users will be instrumental to achieving this goal.

22 American Bar Association, *Perspectives on Predictive Coding and Other Advanced Search Methods for the Legal Practitioner* 58 (2016).

23 As noted in "Notable Developments Since the Convener," *supra* at page 2, in January 2021, the FDA released its "Artificial Intelligence/Machine Learning (AI/ML)-Based Software as a Medical Device (SaMD) Action Plan" (the Action Plan), following up on its earlier version of a "Predetermined Change Control Plan" in premarket submissions. The Action Plan describes the FDA's intention to:

1. Update the proposed framework for AI/ML-based SaMD, including through issuance of Draft Guidance on the Predetermined Change Control Plan;
2. Encourage harmonization of Good Machine Learning Practice development;
3. Conduct a public workshop on how device labeling supports transparency to users and enhances trust in AI/ML-based devices;
4. Support regulatory science efforts to develop methodology for the evaluation and improvement of machine learning algorithms, including for the identification and elimination of bias, and for the evaluation and promotion of algorithm robustness; and
5. Work with stakeholders who are piloting the Real-World Performance process for AI/ML-based SaMD.

The Action Plan is available at <https://www.fda.gov/media/145022/download>.

Regulatory Gaps:

One potential gap for the FDA derives from AI's function as both a product and an evolving system. In other words, pre-market evaluation of an AI in one state of functioning relative to one known use case may be insufficient vetting as AI will evolve both in functioning and application over time. The FDA's Action Plan indicates it recognizes these challenges, despite having traditionally viewed its regulatory authority as extending only to the practice of medicine.

In addition, the vast majority of AI tools are now being developed in a modern cottage industry-type environment, exclusively for in-house use, rather than more familiar and strenuously regulated large commercial deployments. Those more individualized, exclusive-use products will never be vetted by the FDA, casting further doubt on how deeply involved the agency will be in future AI regulation, particularly on issues of liability.

On the one side, the FDA's oversight could be extremely useful in developing trust and transparency around AI; on the other hand, a huge amount of AI will likely escape FDA scrutiny.

FDA Regulation and Liability:

The FDA's current regulatory approach is contributing to less shielding of liability for medical device developers and manufacturers. Typically, medical device developers who go through Class 3 premarket approval are immunized from many forms of state tort law liability. To the extent that the FDA's process is becoming easier and more flexible, with greater emphasis placed on post-launch marketing rather than on the more rigorous upfront approval process, it is expected to offer much lower shields against liability in the future.

The Cures Act²⁴ holds that when developers provide physicians with sufficient information to enable them to understand the AI's recommendation, the associated algorithm or clinical decision support software is not a medical device and need not go through FDA regulation. The portrayal of the developer as a mere provider of information potentially makes physicians "liability crumple zones," as they are deemed to be the ultimate decision makers with sufficient knowledge to justify shifting responsibility for AI choices to them exclusively. How often this will be accurate is questionable, as providers might actually lack sufficient knowledge to communicate risks and benefits of certain AI-enabled technology to patients.

The FTC:

As previously noted, the FTC exerts its broad authority under Section 5 of the FTC Act to protect consumers from unfair and deceptive acts and practices to address general issues regarding data protection, including by extension those that impact the development and use of AI. It takes a three-pronged approach to data protection, focusing on enforcement, policy and consumer and business education.

Enforcement:

A deceptive act under Section 5 is considered not only to be lying to a consumer about how their data will be used, but also conduct likely to mislead a reasonably acting consumer. Section 5 can therefore be used as a basis for an enforcement action challenging the accuracy of representations about the degree of protection a company has afforded data, as it was in *In re Henry Schein Practice Solutions, Inc.*²⁵ Schein was ordered to pay \$250,000.00 to the FTC to settle claims that its software failed to provide industry-standard encryption of sensitive patient information and therefore satisfied HIPAA requirements, per its representations to customers, when it actually used a less complex data masking process that failed to meet accepted encryption standards.

²⁴ Cures Act § 3060(a).

²⁵ Final Order, F.T.C. Docket Number C-4575, May 20, 2016.

Policy:

The FTC's policy pronouncements generally take the form of reports, like the one it issued in 2016 on big data.²⁶ The report recommended operators of algorithms evaluate their models by considering the following questions:

- ▶ How representative is the data set?
- ▶ Does the data model account for biases?
- ▶ How accurate are the predictions made based on big data?
- ▶ Does the reliance on big data raise ethical or fairness concerns?

Education:

Finally, the FTC's educational programs consist largely of speeches and blogs aimed at both consumers and businesses and hearings. In 2020, one commissioner's speech specifically addressed AI.²⁷ At various times, the FTC has held hearings on big data and other areas relevant to AI.

Regulatory Gaps:

Where are the gaps in FTC's regulatory scheme? On the AI front, some companies work in the background, so they are neither consumer-facing nor making representations to consumers, which may prevent the FTC from reaching them using its deception authority. While the FTC can challenge the conduct of such firms using its unfairness authority, unfairness claims present obstacles, such as establishing the substantial injury consumers suffered and whether that injury was reasonably avoidable. This suggests that there is a need for expanded authority for the FTC. Furthermore, given that the FTC is a long-established federal agency regulating a space as novel as AI development, companies may try to challenge the FTC's authority.

Other Regulatory Issues:

An aspirational goal for future regulatory efforts includes standardizing testing for AI to ensure data sets are properly and consistently de-identified and trained, and making the same sets, with comparable guardrails in place, available to all users. AI's inherent variability and dynamic nature, however, may make standardization impossible. These impediments to standardization could be overcome by using discrete sets for specific data types with unique nuances, like images, unstructured and structured electronic health information, labs, and for other elements such as medications.

Using extremely large and homogenous data sets also presents an enhanced potential for adversarial attack. It is much easier to launch a system-wide assault on a single data set. Even when multiple data sets are in use, the more that is known about their construction and the greater the access to that knowledge, the easier it is to fool the set, whether intentionally or unintentionally.

Similar guidelines could conceivably apply to interoperability with regard to other health information systems. Although some guidelines currently exist, coordinating them and ensuring compliance with multiple agency rules simultaneously may be challenging.

Standardization for developing and validating testing sets is acknowledged to be desirable both to enable the FDA to do its job and to aid developers, but the devil, as always, is in the details. Nonetheless, designing standardized, representative data sets for health care AI is an industry goal, although it is not clear if it is an attainable one.

²⁶ "Big Data: A Tool for Exclusion or Inclusion? Understanding the Issues" (FTC Report, January 2016, available at www.americanhealthlaw.org/AI).

²⁷ See Remarks of Commissioner Rebecca Kelly Slaughter, Algorithms and Economic Justice (FTC, January 2020, available at www.americanhealthlaw.org/AI).

Liability: Aligning Stakeholders' Ability to Mitigate Risk with the Value Derived from the Development and Deployment of the AI Solution

AI has now extended into areas that have traditionally been the exclusive responsibility of physicians, like diagnosis of disease, recommendation of treatment and prediction of outcomes. As AI is increasingly used to support or perform traditional medical functions, the use of AI implicates liability theories based on negligence, medical malpractice, enterprise liability, and product liability.

When AI is used for clinical decision making, who or which entity should be liable for negative patient outcomes, and in what circumstances? Physicians have been grappling with this issue for some time.

From physicians' point of view, the person or entity who is most knowledgeable about the implications of AI under the circumstances and is most capable of mitigating the risk of bad outcomes should be liable for them. Aligning the assignment of liability to the actors most capable of risk mitigation creates incentives to mitigate risks associated with AI.

For example, mitigation can occur through product design, development and validation, phases which developers drive and are best positioned to make safe, as well as with the implementation of AI.

Part of the controversy about where liability should reside is based on whether physicians and other end users have enough transparency and understanding when they use a tool to make the appropriate assessments about mitigation. When the tool is autonomous or the end users can't see the underlying data for purposes of accessing risks and benefits or published data validating the tool, those are situations where the developers are best positioned to assume liability. In areas such as use of drugs where physicians are liable for negligence for misuse, physicians have access to extensive published validation about new treatments that describes all aspects of its development, including the details of clinical trials, the composition of the control groups, appropriate dosages, side effects, safety issues, and adverse outcomes. Physicians can then reasonably counsel patients about the risks and benefits associated with the use of the drug. Physicians believe that they should have comparable transparency regarding AI applications.

The issue of transparency into the AI tool is becoming increasingly controversial, however. When a developer claims trade secret protection for a data training set, or imposes gag orders preventing the release of information about known flaws and patient harm, physicians are deprived of the ability to make informed decisions about their use of AI and their ability to explain the risks and benefits to their patients and should not be charged with subsequent liability, from the physicians' standpoint.

Another approach to assigning liability is to determine the extent to which the tool is substituting for a clinician's professional judgment or function. Establishing whether the clinician who employs the AI-enabled tool or the health care IT vendor or developer should be held liable for bad patient outcomes or the degrees of comparative negligence is complicated by the myriad functionalities and the blurring of boundaries between the tasks performed by each.

The 21st Century Cures Act includes a provision that says that if AI gives enough information to enable a user to understand the AI decision, then the AI is not a medical device. This law could result in localizing liability with the physician. The ongoing duty of physicians to oversee AI-enabled technology used in their medical practices and the question of how much the physician must augment the decision-making ability of the tool with their own medical judgment remains unsettled. AI is particularly risky because errors in medical judgment are likely to be replicated across a greater number of patients as compared to errors committed by a human doctor whose actions are likely to harm only a single patient.

With regard to a standard of care for use, much will depend on when AI becomes accepted as a standard tool, comparable to the routine lab tests that are now so familiar and expected to be run in specific circumstances. At that point, community standards may involve a determination of which tool to choose rather than whether

to use an AI tool, although the quantity and variety of tools likely to be available will make this challenging for physicians. The standard of care may ultimately evolve to a point when a decision *not* to use an AI tool could constitute medical malpractice.

Liability determinations as to physicians or entities will continue to be made on a case-by-case basis for at least the immediate future. Determining the appropriate standard of care is a bi-directional issue, particularly where the tool performs better than the human. In those circumstances, there is a need to “retro-analyze” to evaluate if a physician’s *failing* to use the AI-enabled tool harmed the patient.

Moreover, liability for AI decisions may implicate more than traditional tort and malpractice concepts. It is well-recognized that AI has the real potential to improve access to quality care and quality outcomes, but AI done poorly can also exacerbate health inequities that may generate additional liability issues.

The utility of traditional negligence concepts may diminish as AI algorithms improve and doctors rely on AI more heavily for making diagnostic and treatment decisions. Algorithms can change in unexpected ways, making it difficult to apply traditional concepts of foreseeable harm and human causation. Further, it may be impossible to replicate a negative result, because the AI-enabled technology will have learned and changed since its original use.

One suggestion was to encourage shared responsibility for new technologies. A shared responsibility approach recognizes the complementary roles that providers, developers and vendors play in ensuring the safe and effective technology, usability, interoperability, and security of AI-enabled tools. Both the Institute of Medicine²⁸ and the National Academies²⁹ have endorsed the concept of a shared responsibility approach in the context of building health IT to improve patient safety and clinical diagnosis.

Another suggestion was conduct impact assessments to identify sensitive issues that require additional scrutiny and to implement guidelines to inform designers and data scientists.

Data Bias:

Recognizing that a single representative data set might never be achieved, what are the implications for liability for data bias where the data sets may not be representative of the populations they are used on or embody certain historical biases?

First, different levels of data bias exist. From a practical perspective, the challenge is to recognize those levels of bias and identify where bias exists, and put in the best controls to eliminate it.

Bias can exist in the data itself, particularly when the data is historic. Inherent bias can arise from the data set being less than comprehensive and diverse with respect to demographics, like race and gender.

It can also arise from the selection of data, particularly if criteria are applied in such a way as to exclude a particular cohort from inclusion in a clinical trial. Human monitoring, data governance and control, and establishing regular checkpoints can help control for these undesirable results associated with data bias.

Some unresolved questions about data bias include:

- Whether existing theories of anti-discrimination apply when AI produces disparate outcomes?
- Do AI-generated disparate outcomes constitute intentional discrimination?
- What degree of scienter is necessary in terms of whether bias should have been recognized and whether a remedial act was, was not, or should have been taken?

²⁸ Institute of Medicine. 2012. *Health IT and Patient Safety: Building Safer Systems for Better Care*. Washington, DC: The National Academies Press. <https://doi.org/10.17226/13269>.

²⁹ National Academies of Sciences, Engineering, and Medicine. 2015. *Improving diagnosis in health care*. Washington, DC: The National Academies Press.

It has been suggested that AI shows us that there is nothing special about suspect classes, so you can use AI to control for suspect classes recognized by law: race, sex, disability, age, etc. But AI can produce outcomes that are “off” compared to what is considered optimal through very unintuitive kinds of classes. Rules and safe harbors can be set up to require how clean a set must be. But outcomes depend on how optimization takes place.

In optimization, winners and losers may not be immediately obvious. Is it desirable to allow everyone who loses to have a cause of action?

The answer may lie not in traditional tort liability precepts, but through socialization of risk, much like the vaccine compensation system. Holding the occurrence of a bad outcome that would not have occurred had AI not been involved to be the precondition of recovery might prove more effective than attempting to employ fault-based rules.

If screening out bias is particularly important, a fair and effective approach may be to contractually require representations and warranties related to the diversity of the data.

The evolving liability issues may be arising in part as a result of FDA’s current approach regarding liability shielding. Typically, medical devices are immunized from state tort liability. To the extent that devices go through a more flexible process that focuses on post-marketing monitoring, there is less immunization from liability.

An additional question is how AI affects issues of consent. One threshold inquiry is when patients need to know AI is being used in their care. Although this is an evolving discussion, the answer may depend on whether the AI tool is the primary driver of the decision and how much patients need to know. In order to obtain appropriately informed consent, clinicians will need to develop the language to address concepts like continuous learning with patients.

Contracts and AI: Aligning Stakeholders’ Ability to Control/Mitigate the Risk and Value Derived from the Development and Deployment AI Solutions

The liability framework discussed above also informs contract formation. Some Convener participants suggested that the party having the most control over factors giving rise to risk to patients is in the best position to prevent and mitigate the risk, and thus should shoulder contractual liability related to that risk.

In addition, AI agreements should reflect good data stewardship, whether the data is held by health care providers, health plans, or the entities with whom they collaborate to develop AI tools.

Unique Characteristics of AI Development Contracts:

Part of the value proposition for AI is its ability to change over time as it “learns.” Accordingly, the original parameters of the tool are likely to change from the time training is commenced to when it is released to the market. Of course, the fact that the tool evolves over time makes it important for the parties to consider who “controls” the content of the algorithm and thus should shoulder liability related to its performance.

Protection of Patient Information:

AI agreements often involve large amounts of data in various forms. The Convener participants discussed the need to address a number of provisions in those agreements:

- Representations and warranties should be made that the data was obtained lawfully and in compliance with patient consent, if applicable.
- To the extent de-identified data is provided, the contract should prohibit the recipient from re-identifying the data, and should require the recipient to pass on that contractual restriction to any downstream recipient.

- ▶ Many contracts prohibit combining the data with other data sets, which may increase the likelihood of re-identification.
- ▶ The recipient should agree to indemnify the data source with respect to these prohibitions related to re-identification.
- ▶ Security standards should be set out to ensure the vendor is securing the data appropriately per the contract, per NIST guidelines or other relevant standards, and should require the parties to encrypt the data at rest and in transit.

The contract should also describe methods of data destruction when the agreement concludes. If the data will not be aggregated with data from other data sources, the contract should require return or certified deletion of the data. If the data will be aggregated with data from other data sources, the parties should consider whether the data source can be “tagged” and disaggregated for return or destruction at termination. If that is not possible, the data source should consider requiring continuing contractual protections of the data past termination.

The parties should consider the appropriate location of data: is it prudent to bring the data to the tool, or the tool to the data? For data sources, transmitting the data to an AI development partner increases the need to have specific downstream controls for data recipients; doing research on data resources controlled by the data source can alleviate that pressure significantly. On the other hand, commercial partners may be reluctant to bring their tools to the data because of proprietary concerns around the AI algorithm. There may also be limitations in transferring the tools to the data source, such as licensing or technology limitations. Where the data resides and what downstream restrictions are appropriate to place on commercial partners are always hotly debated in these arrangements.

The parties may choose to address data ownership. One perspective is that data ownership may not be as significant as the scope of licenses granted, and determining who can do what with the data, where, and for what purpose.

If the data will be de-identified (which is typical), the agreement should address who is responsible for de-identifying the data, what methodology or standard is being used, the expiration of the de-identification opinion (if the expert methodology is used), and who will bear the costs of de-identification.

To the extent the parties will be collaborating on developing the AI, the parties obviously must address IP ownership and responsibilities for securing patent rights and other commercialization rights.

The parties should address the secondary use of a tool in a context different from the one in which it was originally conceived. This could be covered in the scope of the data use license and by collaboratively setting regular checkpoints throughout the development process to determine when future adjustments to the license are required. Having checkpoints that reveal a potential second use case enables the parties to confer about unexpected results that AI produces, and to update the license and revisit the sufficiency of the data set to avoid health care inequities, should that secondary use be formally explored.

The parties should address approval by an Institutional Review Board (IRB),³⁰ including which party is responsible for obtaining IRB review and communication with the IRB about outcomes and potential adverse events when AI is tested.

³⁰ 21 C.F.R. § 56.102 (g)

Intellectual Property Issues to Address in AI Development Contracts:

The United States Patent and Trademark Office (USPTO) released a report on public views of IP and AI in October 2020 (the Report).³¹ Significantly, it offers no clear definition of AI, which is clearly problematic for building a trusted framework for its use.

The Report is not expected to have much impact on AI at all. The Report discusses at length potentially treating AI as an author or inventor in the IP context, characterizations that have not been widely embraced.

The Report also seems to suggest that nothing novel or different is needed to address AI because the traditional IP framework should afford it sufficient protection. Patents in the medical AI space are therefore likely be fairly limited because of subject matter limitations on patentable subject matter eligibility.³² Enablement and patent breadth are likely to be problematic for AI.

In health care AI specifically, one perspective is that patentability will be less important than trade secrets protections with respect to data and algorithms. Secrecy is problematic for validation, addressing data bias, and obtaining big data sets as opposed to keeping data siloed, although it may ultimately be better for privacy concerns.

Despite the Report's purported summary of public reactions to AI and policy, it has been suggested that the findings do not reflect what's actually going on with the "boots on ground." It fails to address joint patents for the development of AI, where the merging of data and subsequent identification of original ownership of product segments remain challenging. The comments seem to be limited to AI that mimics medical devices, and do not address the complexities that derive from working with unsupervised learning. Nor does it address complications arising from international transactions or attempt to harmonize international rules of patents with those of the United States.

Developers want patent protections for at least portions of the development process or the members of the team that writes the code or other elements. Perhaps subsequent USPTO reports will address the elements of the output of the AI tool to which it is difficult to assign ownership.

How should or could AI or computer-initiated inventions be protected from infringement or copying without acting as a restraint on trade? Tech companies' reliance on open source algorithms makes them less likely to be concerned about IP protections. In fact, IP emerges as a potential structural deterrent to competition in the AI space. More likely sources of protection are traditional non-compete agreements on resourcing, restrictions on licenses, and restrictions on future uses.

As previously discussed, secrecy is the biggest source of protection for AI inventions. Copyright is particularly ill-suited to protecting algorithms or AI products themselves, because of their dynamic evolution, which makes the ability to copyright them as they existed at one point in time of little use. But training on others' copyrighted data raises the question of whether or not the training process infringes the right. Neither data nor facts can be copyrighted, but there are lots of ways of processing data that could be copyrighted. Copyright protection is thus thin.

Conclusions: Primary Elements of a Trusted Framework for Developing AI in Health Care

The specifics of a trusted framework for AI in health care emerge from consideration of various broad and overarching concerns, many of which have not yet been fully explored or resolved. While the Convener did not have time to develop a consensus position, the discussion of the Convener participants pointed to a few

31 "Public Views on Artificial Intelligence and Intellectual Property Policy," USPTO (October 2020, available at www.uspto.gov/sites/default/files/documents/USPTO_AI-Report_2020-10-07.pdf).

32 35 U.S.C. § 101.

important points to develop a trusted framework for the AI in health care: **improving patient experience, serving health equity interests, and achieving better health outcomes.** Accomplishing those ends will require tools that have the following characteristics:

- ▶ The AI tools must be clinically validated;
- ▶ They must improve experiences for both patients and providers;
- ▶ They must be integrated into the physician's workflow;
- ▶ They should address value and costs; and
- ▶ They must be implemented with an eye toward health equity, to prevent this important technology from becoming the exclusive province of well-insured patients treated by sophisticated providers.

The Convener participants also discussed that **a trusted framework strikes an appropriate balance between and among managing legal and regulatory compliance risks, maintaining flexibility to adapt to the dynamic nature of AI innovation, and fueling innovation.** It should embody a continuous, embedded quality improvement process, to prevent the model from becoming static. Ensuring the AI solution's compliance with legal and regulatory requirements will require ongoing testing and re-validation of safety, efficacy and privacy and security protections.

Transparency is important to assure consumers and providers about the tool and whether it is appropriate to use. Transparency with patients complements the adoption of a patient-centered approach to patient care.

Privacy is key in the development, evaluation and implementation of AI. In a related issue, **data security measures** are essential to ensure that unauthorized parties do not have access to patient information used to develop, evaluate and implement AI.

Finally, some Convener participants emphasize that the **regulatory regime should be based on shared accountability**, in which the degree of risk the AI solution aligns with the stakeholders who control risk and derive value from the AI initiative. Measures should be in place for on-going evaluation of AI to ensure that the tool is delivered in the safest and most secure way to patients, being mindful that AI has the potential to replicate errors across large patient sets. The goal is not over-regulation, but recognizing and responding to the needs of an expanded population of stakeholders. The insurance framework must also evolve to recognize the need for shared liability in AI.

Integrating community voices and views, particularly in the area of research on AI, will also contribute to trust in AI. Establishing councils that include patient representatives is one action that can be taken to achieve this result. One perspective is that guardrails should be erected around whether informed consent is required, including how granular the consent should be, the duration of the consent, and the right to revoke consent.

Collaboration from the entire stakeholder community is key. Principles matter but **moving principles to practice** will be the ultimate path to fostering trust. Doing so will require that a broad range of stakeholders have input into both the principles and the practice.

Regulatory flexibility will be crucial to prevent innovation from being stifled. Currently, a disparate system of laws exists that touch on AI but only peripherally; that system needs harmonizing. The creation of a body that monitors privacy and oversees AI, particularly in the health care arena because of the unique risks it presents, could prove helpful and is an avenue that should be explored.

This publication is designed to provide accurate and authoritative information in regard to the subject matter covered. It is provided with the understanding that the publisher is not engaged in rendering legal or other professional services. If legal advice or other expert assistance is required, the services of a competent professional person should be sought.

—From a declaration of the American Bar Association

© 2021 by American Health Law Association.

All rights reserved. No part of this publication may be reproduced, stored in a retrieval system, or transmitted in any form or by any means, electronic, mechanical, photocopying, recording or otherwise, without the express written permission of the publisher.

This White Paper is brought to you by AHLA members and donors like you.

For more information, please visit americanhealthlaw.org.

AHLA's Commitment to Inclusion, Diversity, Equity, and Accessibility

In principle and in practice, the American Health Law Association (AHLA) values and seeks to advance and promote diverse, equitable, inclusive, and accessible participation within the Association for all staff and members.

Guided by these values, the Association strongly encourages and embraces meaningful participation of diverse individuals as it leads health law to excellence through education, information, and dialogue.



American Health Law Association

1099 14th Street, NW, Suite 925

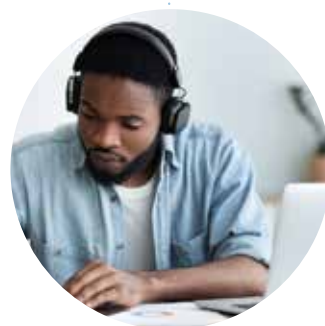
Washington, DC 20005

(202) 833-1100

www.americanhealthlaw.org

A Home for All Health Law Professionals

ACADEMICIANS | CPAs | COMPLIANCE OFFICERS | CORPORATE ATTORNEYS |
HEALTH CARE LAWYERS | FOOD & DRUG LAWYERS | HEALTH CARE CONSULTANTS |
HOSPITAL & NURSING HOME ADMINISTRATORS | HUMAN SERVICES ATTORNEYS |
IN-HOUSE COUNSEL | PHYSICIANS | PUBLIC HEALTH OFFICIALS |
REGULATORY PROFESSIONALS | SOLO PRACTITIONERS | STUDENTS



AHLA has grown to include nearly 13,000 members and over 25,000 engaged health law professionals. From attorneys to compliance professionals, in-house counsel to finance and privacy officers, health care consultants to regulators, all health law professionals interested in health care legal and regulatory issues turn to AHLA to stay up-to-date on the changing health care legal environment.

We are the premier association for all who are interested in health law and we invite you to join AHLA and gain access to a variety of resources:

Continuing Education

Learn from experts in the field and obtain your CLE, CPE, and CCB credits through our educational programming, distance learning and trainings.

Practice Groups and Communities

Connect with a health law community that provides you with knowledge, enables you to leverage and share your expertise, and grow professionally.

News & Analysis

Publications providing in-depth reporting on the latest developments affecting health law.

Mentoring Program and Career Center

No matter where you are in your career, we are here to help match you with a mentor or help find your next step.

Health Law Archive

Rich with past content from across our Practice Groups, in-person program papers, journal issues, newsletters, toolkits, briefings, alerts, and much more.

Personalized Membership Experience

Three membership levels from which to choose, providing you with the resources, savings, and value that best meets your professional needs.

Turn to AHLA for all your health law needs.
Join today at www.americanhealthlaw.org/join.

