

Medical Frontiers in AI Liability

Danny Tobey and Allie Cohen, DLA Piper

In a world of interconnected, thinking machines, patients face new prospects for health and wellness, while providers have the blessing, and curse, of assistive and autonomous new tools for healing the sick. The blessing comes from above-human abilities to detect, diagnose, and treat. The curse comes from new forms of liability for health care providers, and a new set of concepts and cautions that didn't exist in the health care systems of the past.

This article analyzes some of the emerging liability issues at key points where systems interact and responsibilities are divided or unclear—at the margins of machine-provider, machine-patient, and machine-machine interactions. The analysis is not exhaustive (and how could it be given simultaneously evolving standards and concepts for Artificial Intelligence (AI) in health care?); rather, this article highlights some of the unifying concepts, concerns, and cautions the once and future health care provider should be thinking about while embracing these new and promising technologies.

Where Providers and Machines Meet: Who Decides and Who Overrides

What does the health care provider do, after years of academic and clinical training, when experience and insight say one thing and the machine-learning AI says another? Is the machine malfunctioning or is it seeing further? How is the doctor to know, and is she liable for overriding or not overriding?

Continuously learning, uninterpretable algorithms increasingly have the decisional autonomy—but neither the personal agency nor pass-through transparency—that make a legal system based on fault and deterrence an effective regulatory tool. Yet already they can outperform human doctors in several ways, and already doctors are hard pressed to detect false positives and negatives—overriding AI is more complicated than overriding a lab test that does not fit the totality of clinical data. And for those looking to sit out the scrum, there is another complexity—soon, with AI raising the achievable standard of care for patients, failure to adopt may provide new sources of liability.

Several agencies and associations are looking into these questions, attempting to provide frameworks in a rapidly changing environment. Among its numerous guidance documents on software issues following the enactment of the 21st Century Cures Act, the Food and Drug Administration (FDA) has proposed iterations of draft guidance to assess which clinical decision support software qualifies as a “device” and thus is subject to FDA’s oversight and regulation. Some of the FDA’s criteria may be relevant to providers and the assessment of their own liability in adopting these systems. A key aspect of FDA’s analysis begins with explainability—can a doctor independently review the basis for a machine’s recommendation? If so, this likely has implications for the distribution of liability as

well: if doctors can independently review the basis of the AI’s recommendation, responsibility more squarely rests upon them to make the final decision and exercise clinical judgment based on the totality of available data. But who assesses whether a doctor can “independently review the basis” for a machine’s recommendation? If, in defense of malpractice claims, doctors argue they could not, that issue may have to be litigated, to determine whether a physician of ordinary competence could have done so.

And what, exactly, does it mean to “independently review the basis” of an AI’s recommendation? Must the review include the same analysis the algorithm used? The language in FDA’s draft guidance, taken from the 21st Century Cures Act, suggests so, speaking of understanding “the” basis of a machine’s recommendation, not “a” basis or “any” basis that might support it.¹ Tracking the machine’s basis is often not possible—and yet physicians might be able to understand alternative or simpler explanations from the algorithm, allowing them to balance the machine’s “reasoning” against other clinical factors that might have been excluded. Would that process be enough to place responsibility on the health care provider? FDA further speaks of whether the AI is “intended” for health care professionals to rely on primarily, but liability may turn on reality of practice, not intention. Absent *a priori* standards, litigation of these issues will be fact-intensive and complex.

Meanwhile, developers and physicians still must grapple with the learned intermediary doctrine, which has the potential to protect an AI’s manufacturer from failure-to-warn claims for risks they have disclosed to intermediaries, like physicians. A successful learned intermediary defense by a manufacturer can turn into a medical malpractice case for a provider. What this means for providers is that adopting AI means understanding AI, enough to understand and explain its limitations. Some of these concepts will be familiar and portable from other diagnostic tools—standards of care; sensitivity and specificity (false positives and negatives, Area Under the Receiver Operating Characteristic curve). Others may not be: locked versus continuously learning algorithms; disparities between training data and local population data; levels of autonomous versus assistive guidance; algorithmic bias; and, eventually, general versus narrow AI. But, to avoid liability, providers must learn these concepts and, just as importantly, be able to explain them to their patients well enough to meet the duty of care.

More granular than the question of control is the look and feel at the very margins of human-machine interaction. The mechanics of “takeover” (as distinct from the *decision* to take over) is an additional pressure point in doctor-machine interactions. In robotic surgeries, the same code making recommendations to the human operator of a machine can also modify the tactile feedback received from the device: such attempts to



provide more “realistic” input to the operator can also impact the ability to judge the machine’s recommendations. This simulated feedback can also make it physically difficult to perform a takeover. Both makers and consumers of AI should be asking these questions on the front end: how do I takeover, is there an unbiased information stream for doing so, and how easy is it to do? Other liability issues on the provider-machine horizon include “alert fatigue”—the human response to high-sensitivity, low-specificity systems—and “skill atrophy”—the human response to infrequent takeover and override scenarios.

There is also the question of real-world evidence. How a physician interacts with AI in the real world can be very different from how AI works in isolation, in lab tests, or even in clinical studies. Examples exist of approved computer-detection systems that showed superior performance in the lab but sub-par performance in the field: for makers and adopters, how you *implement* human-machine interactions in the real world is an important issue, as is *monitoring* real world performance—reference solely to past pre-market results may miss the mark, and fault in human-machine interactions may reside with makers, adopters, or both.²

FDA has identified other key drivers of risk assessment in medical AI, building on the International Medical Device Regulators Forum’s risk classification system for medical devices.³ The “state of health care situation or condition” being addressed is one, ranging from non-serious to serious to critical. The “significance of information” provided by the software to “the health care decision” is another: does it “inform clinical management,” “drive clinical management,” or “treat or diagnose” a condition? Simply put, the question in these draft guidelines becomes: how important is the machine’s advice and for how serious a condition?

Continuously learning, uninterpretable algorithms increasingly have the decisional autonomy—but neither the personal agency nor pass-through transparency—that make a legal system based on fault and deterrence an effective regulatory tool.

In tandem, the American Medical Association (AMA) is advocating a policy approach in which AI developers are held to liability standards based on transparency (where use of non-disclosure agreements triggers additional pass-through liability to makers) and autonomy (a binary *autonomous versus assistive* approach that spans FDA’s intermingled usage of explainability and inform/drive/treat).⁴ AMA also advocates that “where a mandated use of AI systems prevents mitigation of risk and harm, the individual or entity issuing the mandate must be assigned all applicable liability.”⁵ It remains to be seen if and how these principles would be implemented, as some aspects would require legislative or regulatory action. Meanwhile the common law will continue to evolve. To date, most cases on semi-autonomous and autonomous technologies have been settled quickly, but evolutions in published cases are coming.⁶

Designing these human-machine interactions means balancing risks, both AI and human, because neither is perfect. AI, like much technology, offers tremendous potential: it can

Designing these human-machine interactions means balancing risks, both AI and human, because neither is perfect.

see what people cannot, and often earlier. It never gets tired or bleary at the end of a shift. And yet its mistakes are at scale and often are harder to prove or explain. Careful balancing, thought on the front end, AI literacy, and documented procedures and policies are key. So are AI-specific, rather than boilerplate, terms of service, limitations of liability, representations and warranties, and other negotiated clauses that provide clear risk and responsibility allocation on the front end.

Where Patients and Machines Meet: AI-Telemedicine and the Doctor as Retrospectroscope

Telemedicine commonly follows the standard provider-patient model but uses the internet instead of the office as the site of connection. While remote visits are a huge step forward in expanding access to care, provider availability is still a limiting factor. For that reason, automated telemedicine—which allows more patients to interface directly with medical AI to receive diagnoses, prescriptions, and follow-ups—can be quite attractive. Indeed, patient-machine interfacing is already a part of the menu of telemedicine arrangements. Some services allow patients to get prescriptions online and don't involve physicians until after an algorithm recommends a drug. The efficiency is undeniable—but so is the potential liability. When a physician doesn't interact with the patient until *after* a prescription is written, what does that mean for informed consent and the learned intermediary doctrine? How much should the physician backstop the algorithm, and how much can she do so before the efficiency is lost?

Automated telemedicine brings a new and challenging possibility of mass informed-consent liability. In some instances of the current online prescription model, the only information physicians receive before writing a prescription is (1) a standard form filled out by the patient and (2) a recommendation from an algorithm.⁷ The patient is not necessarily counseled on treatment options and may have no opportunity to ask questions before being told what to take. The patient almost invariably consents to treatment, but the “informed” piece hangs in the balance.

Informed consent is a kind of medical malpractice liability. Under this doctrine, providers must give their patients enough information to allow them to make an intelligent decision.⁸ But with automated telemedicine, treatment plans may be presented in absolutes without a discussion of options or pros and cons—leaving the physician to provide this information. To avoid liability, providers should use their post-hoc patient interaction to explain (1) other available options and (2) their willingness to change the prescription (within medical reason). If patients know that their prescription is simply an option, providers are closer to knowing they have fulfilled their duty.

On the supplier side, automated telemedicine poses a greater risk. Life sciences companies may satisfy their common law duty to warn using the learned intermediary doctrine.⁹ The doctrine is predicated on the idea that providers are in a better position to warn patients of a product's dangers than manufacturers are. But when physicians don't interact with patients until *after* the prescription has been written, and then only through online messaging, physician warnings start to resemble labels, and the barrier of the “intermediary” thins.

Moreover, as AI systems evolve over the next decade, the question of who is the better informed intermediary—doctor or machine—will grow increasingly fraught. It is foreseeable that a judge may hold that the learned intermediary doctrine does not apply when patients are diagnosed and treated using automated telemedicine. FDA has noted that direct-to-patient AI will be an area of focus for its regulatory oversight not just when the software informs management of serious conditions, but when a patient cannot independently review a recommendation's basis for even non-serious conditions.¹⁰ Smart suppliers should be thinking about this now, designing systems and warnings that do not presume the learned intermediary doctrine will remain stable over time.

Where Machines Meet Each Other: Cyber-Physical Systems and Interoperability

The Internet of Things has become a familiar phrase for modern health care lawyers. Now, it's time to add cyber-physical systems (CPS) to the lexicon. CPS are interconnected networks of sensors and smart devices that monitor and interact with their environment. To illustrate, self-driving cars are CPS. They contain sensors that monitor their internal systems and the external environment, allowing them to react in real-time to both physical and digital data. That's why a smart car knows to brake differently when the car in front of it stops short on a sunny day versus a snowy day. Interconnected systems are already part of our health care system and daily lives, and we can expect continued advances in medical CPS (MCPS).

Systems that fall into this category can raise complex FDA regulatory questions, both in terms of the path to market, cybersecurity requirements in seeking FDA clearance or approval, and post-marketing obligations.¹¹ And once products have reached the market, it is not only FDA regulation in play. The Federal Trade Commission (FTC) may take action against companies marketing technologies purporting to offer medical, health, or performance-related benefits, applying high standards for substantiation of claims in these areas.¹²

MCPS are defined in part by interoperability—the ability of two or more software programs or devices to interface. Practitioners and developers alike need to understand the potential liability associated with interoperable devices and who is responsible when patients are harmed by faulty connections. The court in *Hansen v. Baxter Healthcare Corporation* placed the blame on developers, holding that the failing that caused a patient's death was a design defect,¹³ while other scenarios may place the duty on sophisticated health care systems to align their particular menu of IT offerings. The future of interoperability liability is an active subject of debate. In February 2019,

the Department of Health and Human Services announced a proposed rule¹⁴ that would expand the 21st Century Cures Act's ban on information blocking, clearing the path for growth of interoperable devices.¹⁵ Systems such as MCPS present a complex matrix of considerations for device safety and security, and developers should seek legal guidance early and often when adopting these systems.

Conclusion: Frontiers, Walls, and Bridges

The future of medicine is interconnection: machine-provider, machine-patient, and machine-machine. Managing the risks and rewards of these frontline interactions will be a moving target, but best practices are emerging. Seeing around corners is easier for those who keep abreast of cutting-edge developments at the intersection of science, technology, and law. Liabilities can arise anywhere, but these touchpoints *between* humans and machines are a key area where new frictions can be anticipated and managed. **C**

Any views and opinions expressed in this article are those of the authors alone and should not be attributed to AHLA.



Danny Tobey, MD JD, is a physician, successful software entrepreneur, and litigation Partner of DLA Piper, a global law firm in over 90 offices and 40 countries. Danny is a leader of the firm's global AI and data science team. He advised the AMA on its health care AI policies and was honored by the Library of Congress in 2019 for

his work on AI and the law. Danny has assisted leading life sciences and health care companies in matters involving FDA, CDC, CMS, HHS, private payors, and commercial litigants. He is a graduate of Harvard College, UT Southwestern Medical School, and Yale Law School, where he received the Gruter Prize for the best work on biological sciences and the law.



Allie Cohen, JD MBE, is a law clerk in DLA Piper's Product Liability group, working with clients in the health care and life sciences sectors. She has conducted research on developmental genetics in *Drosophila melanogaster* and has peer-reviewed publications in molecular biology and genetics. Allie graduated

from Cornell University as a Rawlings Cornell Presidential Research Scholar, with Honors in Biological Sciences. She received her J.D. from the University of Pennsylvania Carey Law School and her Master of Bioethics from the Perelman School of Medicine.

The authors would like to thank their colleagues Rebecca McKnight, Kristi Kung, and Andy Serwin of DLA Piper's FDA, healthcare, and cybersecurity practices for their input.

[A]s AI systems evolve over the next decade, the question of who is the better informed intermediary—doctor or machine—will grow increasingly fraught.

Endnotes

- 1 H.R. 34, 114th Cong. (2016) (enacted); Clinical Decision Support Software 84 Fed. Reg. § 51167 (proposed Sept. 27, 2019). In other places, the language of the Cures Act suggests otherwise, explaining the "independent basis" standard as "so that it is not the intent that user rely primarily on any such recommendation." But that language blends the threshold prong of explainability with the separate prong of "degree of reliance" on the AI, which FDA raises later in its use of the International Medical Device Regulators Forum's risk matrix for enforcement discretion. An AI can be any combination of explainable/unexplainable and assistive/autonomous. The Cures Act's equating of these two independent factors may cause problems.
- 2 See, e.g., Fenton JJ, Taplin SH, Carney PA, Abraham L, Sickles EA, D'Orsi C, et al. *Influence of computer-aided detection on performance of screening mammography*, 356 NEW ENG. J. MED. 1399-409 (2007).
- 3 Clinical Decision Support Software 84 Fed. Reg. 51167 (proposed Sept. 27, 2019); see also FOOD AND DRUG ADMIN. GLOBAL APPROACH TO SOFTWARE AS A MEDICAL DEVICE (2017), <https://www.fda.gov/medical-devices/software-medical-device-samd/global-approach-software-medical-device-software-medical-device>.
- 4 AMERICAN MED. ASS'N, *Health care AI must boost the quadruple aim to move forward*, June 12, 2019, <https://www.ama-assn.org/practice-management/digital/health-care-ai-must-boost-quadruple-aim-move-forward>.
- 5 *Id.*
- 6 See Danny Tobey, *Software Malpractice in the Age of AI: A Guide for the Wary Tech Company*, Proceedings of the 2018 AAAI/ACM Conference on AI, Ethics, and Society, New Orleans, LA, USA—Feb. 02–03, 2018.
- 7 Yasmin Gagne & Burt Helm, *Buying prescription drugs online is easier than ever. But there are side effects*. FAST COMPANY (Sept. 11, 2019) (explaining "an algorithm analyzes patients' symptoms, then submits a diagnosis for a human doctor to review"), <https://www.fastcompany.com/90439144/machine-learning-could-transform-medicine-should-we-let-it>.
- 8 A successful informed consent claim requires (1) existence of material risk unknown to the patient, (2) failure to disclose the risk, (3) the patient would have made a different choice if the risk was disclosed, and (4) injury. *Canterbury v. Spence*, 464 F.2d 772, 786–87 (D.C. Cir. 1972).
- 9 *Ackermann v. Wyeth Pharmaceuticals*, 526 F.3d 203 (5th Cir. 2008).
- 10 Clinical Decision Support Software 84 Fed. Reg. 51167 (proposed Sept. 27, 2019) at 23.
- 11 See, e.g., FDA's October 2019 Urgent/11 Cybersecurity Vulnerabilities May Introduce Risks During Use of Certain Medical Devices, October 2018 draft guidance on Content of Premarket Submissions for Management of Cybersecurity in Medical Devices, and December 2016 Postmarket Management of Cybersecurity in Medical Devices.
- 12 See, e.g., 80 Fed. Reg. 4575 (Jan. 28, 2015) (Focus Education, LLC, Analysis of Proposed Consent Agreement to Aid Public Comment); 80 Fed. Reg. 11437 (Mar. 3, 2015) (Health Discovery Corporation, Analysis of Proposed Consent Agreement to Aid Public Comment).
- 13 No. 89043, 2002 WL 93106 (Ill. Jan. 25, 2002). The providers settled.
- 14 84 Fed. Reg. 7424 (Mar. 4, 2019) proposed to be codified at 45 C.F.R. § 170-171.
- 15 Pub. L. No. 114-255 (2016). The comment period for the proposed rule is now closed, but will likely be finalized by 2022. Under the Medicare Modernization Act of 2003, Medicare final regulations have three years from a proposed or interim final rule to be published to the *Federal Register* as final rules.