

Witness Testimony to House Committee on Insurance

Good morning, Chairman Oliverson, Vice-Chairwoman Johnson, and representatives on this committee. Thank you for inviting me to participate in this hearing. I appreciate your service to the people of the great state of Texas, and it is an honor to testify before the Committee today on the topic of Artificial Intelligence and health care. My name is Kev Coleman, and I am a Visiting Research Fellow at Paragon Health Institute, a health policy research center for improved outcomes in the public and private sectors. Alongside my work in health care research, I am also a former software technology executive and startup consultant. My testimony today represents my views alone.

Artificial intelligence, or AI as its abbreviated, is not a singular technology. Instead, it represents multiple categories of programming that emulate human learning and reasoning to various degrees. In some cases, these categories evidence significant differences from one another. The four forms of AI that are particularly relevant to health care are:

- Large Language Models
- Machine Learning
- Artificial Neural Networks
- Generative AI

These four types of AI have already made remarkable inroads within the field of medicine and achieved results that, at times, defy the imagination. We now have AI software that can review a single low-radiation chest scan and predict a patient's lung cancer risk for the following six years without input from a radiologist. In some instances, the system has been able to detect early lung cancer signs that radiologists did not recognize until lung nodules were visible on scans years later. While many of AI's medical applications are associated with medical imaging, AI is found throughout health care in areas such as drug discovery, administrative automation, precision medicine, patient care, mental health, population health management, claims fraud detection, and surgical robotics.

The subject of AI is as closely related to data as it is to software algorithms. Unlike traditional software where programmed commands account for much of an application's performance, AI software is reliant on large data sets to train the system to produce desired outputs. Given that many health care AI systems need patient data for their training, there are concerns around consumer privacy, a matter I shall touch on later in my testimony.

Turning from AI data issues, there is the economic context of current AI development. AI's rise within medicine coincides with growing anxiety over the American health care system's fiscal challenges. In 2009, the Social Security Advisory Board warned the nation's health care cost trajectory was "unsustainable" and "perhaps the most significant threat to the

long-term economic security of workers and retirees.”¹ At that time, the United States spent approximately \$2.5 trillion on health care, which represented \$8,160 per U.S. resident.² By 2022, the United States spent \$13,493 per person on health care, totaling \$4.5 trillion.³ This amount represented 17.3 percent of the nation’s gross domestic product.⁴ In comparison, the nation spent only 6.2 percent of its GDP on health care in 1970.⁵

The financial unsustainability of American health care suggests that, in addition to its potential for clinical improvements, AI should also be explored as an instrument for cost reduction. Its ability to replicate or exceed human reasoning opens the possibility of replacing some high-cost human labor with lower cost AI functionality.

This brings us now to the subject of AI regulation. Among the complications for policymakers considering AI regulation is the technology’s integration across the health care spectrum. Notably, it will be as prevalent outside clinical settings as it is inside. Email systems will have AI for SPAM filtering, threat detection, composition assistance, and other functionality. The same will be true for accounting software, H.R. software, word processing, spreadsheet, customer relationship management, and search engine marketing. Poorly reasoned and overly broad regulation of health care AI could cast wide net that inflates compliance costs while achieving little to improve patient safety.

Equally problematic is the overlap between traditional software functionality and AI functionality. Traditional software, for example, can make statistical predictions and recommendations, detect patterns, generate novel outputs that were not pre-programmed, and receive inputs based in natural language. Once again, broad regulation intended for AI would bring chaos to simpler software systems whose functionality would trigger the same regulatory obligations.

Policymakers drafting laws related to AI must be exceptionally careful to specify criteria to differentiate AI from non-AI systems. Furthermore, given the differences in operation and risk even among categories of AI, a blanket AI regulation may be irrelevant for some AI implementations and counterproductive for others. This reality necessitates a cautious approach to regulation that leans heavily on the expertise of independent AI experts. By

¹ Social Security Advisory Board, The Unsustainable Cost of Health Care, September 8, 2009, https://s3-us-gov-west-1.amazonaws.com/cg-778536a2-e58c-44f1-9173-29749804ec54/uploads/2023/12/2009-TheUnsustainableCostofHealthCare_2009.pdf

² Kaiser Family Foundation, "Trends in Health Care Costs and Spending," March 2009, https://www.kff.org/wp-content/uploads/2013/01/7692_02.pdf

³ See the Centers for Medicare and Medicaid Services’ historical national health expenditure data at <https://www.cms.gov/data-research/statistics-trends-and-reports/national-health-expenditure-data/historical>

⁴ Ibid.

⁵ Emma Wager et al., "How Does Health Spending in the U.S. Compare to Other Countries?" Peterson-KFF Health System Tracker, (January 23, 2024), <https://www.healthsystemtracker.org/chart-collection/health-spending-u-s-compare-countries/>

independent, I mean persons not in the employ of large AI development entities with deep pockets. Such large companies often prefer higher regulation because it can reduce competition entering the market.

The mention of large developers raises another consideration in the potential regulation of health care AI: the need for a framework whose compliance costs will not bankrupt innovative start-ups that lack the financial resources of a Microsoft or Oracle. If we fail to provide such a framework, the market may be reduced to a few powerful companies controlling one of the most important technologies of the twenty-first century. The risk of excessive market consolidation should motivate federal regulators to examine acquisitions carefully from an antitrust perspective. Time Magazine has noted that “Apple, Microsoft, Google, Meta, and Amazon have collectively acquired at least 89 AI companies over the last decade, and those acquisitions tended to target younger startups, a signal that the tech giants may be targeting innovative AI firms before they pose a competitive threat.”⁶ For AI start-ups, monopolies are not only a concern with respect to software applications but also the infrastructure resources AI software needs for its massive data processing.

Alongside these considerations, regulators must also resist unrealistic and unhelpful standards of perfection for AI medical devices. Human doctors are not perfect, and we still allow their practice. With respect to AI safety, we should insist on an accuracy rate at least the same as exhibited by clinicians. Perfection is not always possible in the context of medicine and the expectation would prevent lower cost AI solutions from replacing expensive human labor that has been at the heart of the nation’s health care cost problem. If an AI system can be demonstrated to have the accuracy equivalent to human doctors, safety risks are not being increased.

The risk of harm is a key dimension in this technology’s appropriate regulation. Our greatest attention and granularity of response should be directed AI implementations where the risk of patient harm is highest. Likewise, where there is no risk of patient harm (for example, using AI to detect claims fraud in Medicaid), the impulse to regulate should lessen.

The matter of acceptable risk is especially pertinent to autonomous AI solutions where system accuracy and low patient risk eliminate the necessity for clinician assistance. For such tools it is crucial that clinical assistance not be made a regulatory requirement if an AI system can empirically demonstrate to regulators that it satisfies the following three criteria:

- Accuracy levels equal to or exceeding the average rate for clinicians performing the same function
- No amplification of health risks as compared to when the same function is performed by a clinician

⁶ Jack Corrigan, "What Google’s Antitrust Defeat Means for AI." Time. (August 29, 2024), <https://time.com/7015493/google-antitrust-defeat-ai-monopolies/>

- Output communications that are comprehensible and actionable for the patient

Given the continuing evolution of AI software systems, it is also advisable that federal regulators provide an economical pathway for innovators to re-apply for FDA approval on their devices where the functionality remains the same, but system autonomy increases over time. On this front, the FDA could leverage work already performed by the U.S. Department of Transportation for self-driving vehicles with differing levels of system autonomy (e.g. driver-assistance vs self-driving).⁷

AI regulators also must become comfortable with the new levels of complexity. An artificial neural network, for example, can have millions of neurons contributing to its output. As a consequence, the system can produce desired clinical results through its response to data correlations developers cannot easily recognize and tease out of the system. In some cases, this may mean the system produces clinically desirable results although the process may not be fully understood even by system developers. While this may seem controversial, this situation already exists in pharmacology, and has for decades, where we have FDA-approved medicines where the scientific basis for their positive results are not always understood.⁸ If we can live with this for pharmacology, I believe we can do the same with AI so long as the AI system can produce empirically verifiable and repeatable outcomes acceptable to the safety standards of the FDA. These safety standards may also require more scrutiny of the underlying data used to train the AI system.

In conclusion, I hope these brief testimony illuminated challenges around AI regulation and encourage a thoughtful and cautious approach that prioritizes regulation around concrete patient risk rather than the hype and fearmongering that too often surrounds discussions of artificial intelligence.

Thank you for the opportunity to testify before the Committee today, and I look forward to your questions.

⁷ National Highway Traffic Safety Administration, "Standing General Order on Crash Reporting," <https://www.nhtsa.gov/laws-regulations/standing-general-order-crash-reporting>

⁸ Carolyn Y. Johnson, "One big myth about medicine: We know how drugs work." Washington Post. (July 23, 2015), <https://www.washingtonpost.com/news/work/wp/2015/07/23/one-big-myth-about-medicine-we-know-how-drugs-work/>