

June 12, 2023

The Honorable Alan Davidson
Assistant Secretary of Commerce for Communications and Information
National Telecommunications and Information Administration
U.S. Department of Commerce
1401 Constitution Ave NW
Washington, DC 20230

**RE: Comments of the Connected Health Initiative to the National
Telecommunications and Information Administration on its AI
Accountability Policy Request for Comment (NTIA–2023–0005)**

I. Introduction & Statement of Interest

ACT | The App Association (App Association) appreciates the opportunity to provide input to the National Telecommunications and Information Administration (NTIA) on developing a productive artificial intelligence (AI) accountability ecosystem.¹

The App Association is a global trade association for small and medium-sized technology companies. Our members are entrepreneurs, innovators, and independent developers within the global app ecosystem that engage with verticals across every industry. We work with and for our members to promote a policy environment that rewards and inspires innovation while providing resources that help them raise capital, create jobs, and continue to build incredible technology. Today, the value of the ecosystem the App Association represents—which we call the app economy—is approximately \$1.7 trillion and is responsible for 5.9 million American jobs, while serving as a key driver of the \$8 trillion internet of things (IoT) revolution.² Alongside the world's rapid embrace of mobile technology, our members create the innovative solutions that utilize AI to power IoT across various modalities and segments of the economy.

AI is an evolving constellation of technologies that enable computers to simulate elements of human thinking, such as learning and reasoning. An encompassing term, AI entails a range of approaches and technologies, such as machine learning (ML), where algorithms use data, learn from it, and apply their newly-learned lessons to make informed decisions, and deep learning, where an algorithm based on the way neurons and synapses in the brain change as they are exposed to new inputs allows for

¹ <https://www.federalregister.gov/documents/2023/04/13/2023-07776/ai-accountability-policy-request-for-comment>

² ACT | The App Association, State of the U.S. App Economy: 2020 (7th Edition) (Apr. 2020), available at <https://actonline.org/wp-content/uploads/2020-App-economy-Report.pdf>

independent or assisted decision-making. AI-driven tools are having, and will continue to have, substantial direct and indirect effects on Americans. Some forms of AI are already being used to improve American consumers' lives today – for example, AI is used to detect financial and identity theft and to protect the communications networks upon which Americans rely against cybersecurity threats. Moving across use cases and sectors, AI has incredible potential to enable faster and better-informed decision making through cutting-edge distributed cloud computing. For example, healthcare treatments and patient outcomes stand poised to improve disease prevention and conditions, as well as efficiently and effectively treat diseases through automated analysis of x-rays and other medical imaging. From a governance perspective, AI solutions will derive greater insights from infrastructure and support efficient budgeting decisions. It is estimated that AI technological breakthroughs will represent a \$126 billion market by 2025.³

Even now, consumers are encountering AI in their lives incrementally through the improvements they have seen in computer-based services they use, typically in the form of streamlined processes, image analysis, and voice recognition, all forms of what we consider “narrow” AI. These narrow applications of AI already provide great societal benefit. As AI systems, powered by streams of data and advanced algorithms, continue to improve services and generate new business models, the fundamental transformation of economies across the globe will only accelerate. Nonetheless, AI's growing use raises a variety of challenges, and some new and unique considerations, for policymakers as well as those making AI operational today. The App Association appreciates NTIA's exploration of policies to provide reliable guidance to stakeholders to reassure end-users that AI systems are legal, effective, ethical, safe, and otherwise trustworthy.

The App Association has worked proactively to develop consensus around AI governance and policy questions from across its diverse and innovative community of small businesses. As a result of these consensus-building efforts, the App Association has created comprehensive policy principles for AI governance,⁴ which we append to this comment and urge NTIA (and other policymakers) to align with. **Notably, the App Association's policy principles for AI governance and policy address quality assurance and oversight, recommending that any AI policy framework utilize risk-based approaches to ensure that the use of AI aligns with the recognized standards of safety, efficacy, and equity. Our AI policy principles also prioritize ensuring the appropriate distribution and mitigation of risk and liability by providing that those in the value chain with the ability to minimize risks based on their knowledge and ability should have appropriate incentives to do so.**

³ McKinsey Global Institute, *Artificial Intelligence: The Next Digital Frontier?* (June 2017), available at <https://www.mckinsey.com/~media/McKinsey/Industries/Advanced%20Electronics/Our%20Insights/How%20artificial%20intelligence%20can%20deliver%20real%20value%20to%20companies/MGI-Artificial-Intelligence-Discussion-paper.ashx>.

⁴ The App Association's Policy Principles for Artificial Intelligence are included in this comment as **Appendix A**.

II. Specific Input on AI Accountability for NTIA's Consideration

AI accountability at its core refers to the promise to ensure that the AI systems that developers design, build, or deploy function properly and responsibly throughout their lifecycle. This means that throughout the course of AI production, developers must continuously examine their work to understand the limitations of the datasets used and account for and manage efficacy, while mitigating data bias and meeting consumer expectations (e.g., privacy).

While we appreciate NTIA's work to advance AI accountability, the App Association strongly urges for a coordinated effort across both executive and independent agencies. Already, numerous regulatory agencies, some cross-sectoral and others sector-specific, are considering or advancing regulatory proposals that would take starkly different approaches to AI accountability. Some of these proposals are poised to put significant hurdles in place for the development and use of AI through one-size-fits-all approaches that have nominal public benefit at best, such as the Department of Health and Human Services Office of Civil Rights' proposed approach to preventing discriminatory outcomes in healthcare,⁵ which the App Association's dedicated digital health advocacy effort, the Connected Health Initiative,⁶ has detailed its views on publicly (and we encourage NTIA's consideration of these viewpoints as a leading example of sector-specific misalignment with other leading Administration efforts, such as that of the National Institute of Standards and Technology [NIST]⁷). In some cases, such proposals are being developed based on speculative and undemonstrated harms.⁸ NTIA, along with other cross-sectoral subject matter expert agencies in the federal government such as NIST, should take immediate steps to ensure a harmonized and informed approach to AI governance.

Many entities, both public and private, are actively engaging in efforts to create and enforce AI accountability frameworks, which may lead to the creation of trusted audits, assessments, and certifications. While this area continues to evolve, we strongly urge for NTIA's alignment with NIST's efforts to develop a voluntary artificial intelligence risk management framework (AI RMF), which aims to help designers, developers, users, and evaluators of AI systems evolve in knowledge, awareness, and best practices to better manage risks across the AI lifecycle.⁹ NIST's AI RMF is best positioned to guide federal government efforts in addressing AI due to NIST's expertise and its collaborative

⁵ Nondiscrimination in Health Programs and Activities, 87 FR 47824 (Aug. 4, 2022); the App Association's Connected Health Initiative detailed views on this HHS OCR proposal are included in this comment as **Appendix B**.

⁶ See <https://connectedhi.com/>.

⁷ <https://www.nist.gov/itl/ai-risk-management-framework>.

⁸ Trade Regulation Rule on Commercial Surveillance and Data Security, 87 FR 51273 (Aug. 22, 2022); App Association views provided to the Federal Trade Commission in response to its Advanced Notice of Proposed Rulemaking are available at <https://www.regulations.gov/comment/FTC-2022-0053-1089>.

⁹ <https://www.nist.gov/itl/ai-risk-management-framework>.

and open approach to developing the AI RMF, similar to NIST's Cybersecurity Framework.¹⁰ It is in the public's best interest that the NIST AI RMF's scaled, risk-based approach serve as a basis for both executive and independent agencies' approach to AI risk management and governance; and that NTIA take active steps to bring federal agencies into alignment with this approach.

III. Conclusion

The App Association appreciates NTIA's consideration of the above (and appended) views and we urge NTIA to contact the undersigned with any questions or ways that we can assist moving forward.

Sincerely,



Brian Scarpelli
Senior Global Policy Counsel

Leanna Wade
Regulatory Policy Associate

ACT | The App Association
1401 K St NW (Ste 501)
Washington, DC 20005
202-331-2130

¹⁰ <https://www.nist.gov/cyberframework>.

ACT | The App Association's Policy Principles for Artificial Intelligence

Artificial intelligence (AI) is an evolving constellation of technologies that enable computers to simulate elements of human thinking, such as learning and reasoning. An encompassing term, AI entails a range of approaches and technologies, such as machine learning (ML), where algorithms use data, learn from it, and apply their newly-learned lessons to make informed decisions, and deep learning, where an algorithm based on the way neurons and synapses in the brain change as they are exposed to new inputs allows for independent or assisted decision-making. Already, AI-driven algorithmic decision tools and predictive analytics have substantial direct and indirect effects in consumer and enterprise context, and show no signs of slowing in the future.

Across use cases and sectors, AI has incredible potential to improve consumers' lives through faster and better-informed decision-making, enabled by cutting-edge distributed cloud computing. Even now, consumers are encountering AI in their lives incrementally through the improvements they have seen in computer-based services they use, typically in the form of streamlined processes, image analysis, and voice recognition, all forms of what we consider "narrow" AI. These narrow applications of AI already provide great societal benefit. As AI systems, powered by streams of data and advanced algorithms, continue to improve services and generate new business models, the fundamental transformation of economies across the globe will only accelerate.

Nonetheless, AI also has the potential to raise a variety of unique considerations for policymakers. ACT | The App Association appreciates the efforts to develop a policy approach to AI that will bring its benefits to all, balanced with necessary safeguards to protect consumers.

To guide policymakers, we recommend the following principles for action:

1. **AI Strategies:** Many of the policy issues raised below involve significant work and changes that will impact a range of stakeholders. The cultural, workforce training and education, data access, and technology-related changes associated with AI will require strong guidance and coordination. National AI strategies incorporating guidance on the issues below will be vital to achieving the promise that AI offers to consumers and entire economies. We believe it is critical that countries also take this opportunity to encourage civil society organizations and private sector stakeholders to begin similar work.
2. **Research:** Policy frameworks should support and facilitate research and development of AI by prioritizing and providing sufficient funding while also ensuring adequate incentives (e.g., streamlined availability of data to developers, tax credits) are in place to encourage private and non-profit sector research. Transparency research should be a priority and involve collaboration among all affected stakeholders who must responsibly address the ethical, social, economic, and legal implications that may result from AI applications.

3. **Quality Assurance and Oversight:** Policy frameworks should utilize risk-based approaches to ensure that the use of AI aligns with the recognized standards of safety, efficacy, and equity. Providers, technology developers and vendors, and other stakeholders all benefit from understanding the distribution of risk and liability in building, testing, and using AI tools. Policy frameworks addressing liability should ensure the appropriate distribution and mitigation of risk and liability. Specifically, those in the value chain with the ability to minimize risks based on their knowledge and ability to mitigate should have appropriate incentives to do so. Some recommended guidelines include:
 - Ensuring AI is safe, efficacious, and equitable.
 - Supporting that algorithms, datasets, and decisions are auditable.
 - Encouraging AI developers to consistently utilize rigorous procedures and enabling them to document their methods and results.
 - Requiring those developing, offering, or testing AI systems to provide truthful and easy to understand representations regarding intended use and risks that would be reasonably understood by those intended, as well as expected, to use the AI solution.
 - Ensuring that adverse events are timely reported to relevant oversight bodies for appropriate investigation and action.
4. **Thoughtful Design:** Policy frameworks should require design of AI systems that are informed by real-world workflows, human-centered design and usability principles, and end-user needs. AI systems solutions should facilitate a transition to changes in the delivery of goods and services that benefit consumers and businesses. The design, development, and success of AI should leverage collaboration and dialogue among users, AI technology developers, and other stakeholders in order to have all perspectives reflected in AI solutions.
5. **Access and Affordability:** Policy frameworks should ensure AI systems are accessible and affordable. Significant resources may be required to scale systems. Policymakers should take steps to remedy the uneven distribution of resources and access and put policies in place that incent investment in building infrastructure, preparing personnel and training, as well as developing, validating, and maintaining AI systems with an eye toward ensuring value.
6. **Ethics:** The success of AI depends on ethical use. A policy framework will need to promote many of the existing and emerging ethical norms for broader adherence by AI technologists, innovators, computer scientists, and those who use such systems. Policy frameworks should:
 - Ensure that AI solutions align with all relevant ethical obligations, from design to development to use.
 - Encourage the development of new ethical guidelines to address emerging issues with the use of AI, as needed.
 - Maintain consistency with international conventions on human rights.
 - Ensure that AI is inclusive such that AI solutions beneficial to consumers are developed across socioeconomic, age, gender, geographic origin, and other groupings.
 - Reflect that AI tools may reveal extremely sensitive and private information about a user and ensure that laws protect such information from being used to discriminate against certain consumers.

7. **Modernized Privacy and Security Frameworks:** While the types of data items analyzed by AI and other technologies are not new, this analysis will provide greater potential utility of those data items to other individuals, entities, and machines. Thus, there are many new uses for, and ways to analyze, the collected data. This raises privacy issues and questions surrounding consent to use data in a particular way (e.g., research, commercial product/service development). It also offers the potential for more powerful and granular access controls for consumers. Accordingly, any policy framework should address the topics of privacy, consent, and modern technological capabilities as a part of the policy development process. Policy frameworks must be scalable and assure that an individual's data is properly protected, while also allowing the flow of information and responsible evolution of AI. This information is necessary to provide and promote high-quality AI applications. Finally, with proper protections in place, policy frameworks should also promote data access, including open access to appropriate machine-readable public data, development of a culture of securely sharing data with external partners, and explicit communication of allowable use with periodic review of informed consent.
8. **Collaboration and Interoperability:** Policy frameworks should enable eased data access and use through creating a culture of cooperation, trust, and openness among policymakers, AI technology developers and users, and the public.
9. **Bias:** The bias inherent in all data, as well as errors, will remain one of the more pressing issues with AI systems that utilize machine learning techniques in particular. Any regulatory action should address data provenance and bias issues present in the development and uses of AI solutions. Policy frameworks should:
 - Require the identification, disclosure, and mitigation of bias while encouraging access to databases and promoting inclusion and diversity.
 - Ensure that data bias does not cause harm to users or consumers.
10. **Education:** Policy frameworks should support education for the advancement of AI, promote examples that demonstrate the success of AI, and encourage stakeholder engagements to keep frameworks responsive to emerging opportunities and challenges.
 - Consumers should be educated as to the use of AI in the service they are using.
 - Academic education should include curriculum that will advance the understanding of and ability to use AI solutions.

The App Association represents more than 5,000 small business software application development companies and technology firms across the mobile economy. Our members develop innovative applications and products that meet the demands of the rapid adoption of mobile technology and that improve workplace productivity, accelerate academic achievement, monitor health, and support the global digital economy. Our members play a critical role in developing new products across consumer and enterprise use cases, enabling the rise of the internet of things (IoT). Today, the App Association represents an ecosystem valued at approximately \$1.7 trillion that is responsible for millions of jobs around the world.

For more information, please visit www.actonline.org.

ConnectedHealthInitiative

April 14, 2022

The Honorable Xavier Becerra
Secretary
U.S. Department of Health and Human
Services
200 Independence Avenue Southwest
Washington, District of Columbia 20201

Melanie Fontes Rainer
Director
U.S. Department of Health and Human
Services
Office for Civil Rights
200 Independence Avenue Southwest
Washington, District of Columbia 20201

The Connected Health Initiative (CHI) represents a diverse coalition of stakeholders that spans the healthcare and technology sectors, all of whom support the expanded use of connected health technologies in healthcare. We write to provide further consensus recommendations on the Department of Health and Human Services' (HHS) Office of Civil Rights (OCR) proposed rule on Section 1557 of the Affordable Care Act (ACA), specifically its proposal to update § 92.210 to "make explicit that covered entities are prohibited from discriminating through the use of clinical algorithms on the basis of race, color, national origin, sex, age, or disability under Section 1557" without clarification on how to detect and mitigate potential algorithmic bias.¹

Below, we share our concerns with proposals poised to impact the use of artificial intelligence (AI) in healthcare and urge for the withdrawal of the technology-specific proposal addressing covered entities' use of algorithms in its 1557 nondiscrimination rules. Alternatively, should OCR nonetheless choose to retain AI-specific provisions in its new 1557 nondiscrimination rules, we offer several ways that these AI-specific provisions in the 1557 nondiscrimination rules should be revised to protect patient safety and equity while supporting innovation in healthcare AI.

CHI is a not-for-profit that seeks to advance policies that will provide the infrastructure and policy environment to support, as well as incent, the use of cutting-edge digital health products and services in both the prevention and treatment of disease. CHI's AI Task Force has developed a set of health AI policy principles to help guide policymakers to effectively address the role of AI in healthcare.² Evidence demonstrates that digital health improves patient care, reduces hospitalizations, helps avoid complications, and improves patient engagement. Leveraging wide ranges of datasets, including patient-generated health data, with AI tools holds incredible promise for equitably advancing value-based care in research, health administration and operations, population health, practice delivery improvement, and direct clinical care. To achieve this potential, government policies must be put in place to support building infrastructure and preparing and training personnel, as well as developing, testing, validating, and maintaining AI systems to ensure value. AI tools are also critical in meeting the Administration's priorities, such as reducing disparities. CHI shares the Administration's commitment to advancing health equity in this regard.³

¹ 87 FR 47884.

² Connected Health Initiative, Policy Principles for Artificial Intelligence in Health, available at <https://connectedhi.com/wp-content/uploads/2022/02/Policy-Principles-for-AI.pdf>

³ <https://www.whitehouse.gov/briefing-room/presidential-actions/2023/02/16/executive-order-on-further-advancing-racial-equity-and-support-for-underserved-communities-through-the-federal-government/>.

We appreciate HHS' efforts to date to responsibly bring the benefits of AI to patients in a way that advances health equities and benefits all providers and patients. For example, the Centers for Medicare and Medicaid Services have taken a number of important steps to make AI's benefits available to more caregivers and patients, including updating its Medicare Physician Fee Schedule rules to provide national payment rates for AI's responsible use in addressing specific use cases, such as in diabetic retinopathy, and integrating AI into value-based care, specifically in various Quality Payment Program Merit-based Incentive Payment System quality measures.

In its proposed rule, OCR proposes to make explicit that covered entities are prohibited from discriminating, through the use of clinical algorithms, on the basis of race, color, national origin, sex, age, or disability under Section 1557, and requests input on the appropriate scope and specificity of such a requirement. While we share HHS' goal of advancing the use of beneficial algorithms by covered entities, share concerns with potential discriminatory outcomes resulting from the use of health AI tools and services, and support the intent of the 1557 rule as a whole, HHS' proposals targeting AI raise a number of concerns, including:

- HHS' evaluation of various use cases demonstrating its concerns with health AI-related discriminatory outcomes does not adequately differentiate root causes for the outcomes it seeks to avoid.
- HHS' proposal to explicitly address an emerging technology area (AI) does not appear to consider ongoing developments in standardization, and further raises the risk of technology terms and capabilities evolving more quickly than regulations can be updated.
- Our community is working to develop a consensus standard on how to validate that biases are being identified and appropriately mitigated, and to establish an adequate infrastructure of test beds for making such standards operational. For example, providers, technology developers, governments, and others continue to address how to make AI data sets appropriately representative of the populations/communities AI tools are intended to serve and benefit.
- HHS' proposal appears to omit that providers rely on a health AI manufacturer's intended uses, whether the AI meets the definition of a medical device or not, and that its proposal would force covered entities to police their own supply chains for AI tools and services, despite realities that would make such efforts impracticable (for example, it is often infeasible to require a covered entity to audit AI and/or the datasets used to train AI they purchase). Further, the additional steps that covered entities would need to take to comply with HHS' proposed requirement are likely to contribute to providers' already strained workload and further contribute to burnout.
- HHS' proposal does not account for the fact that some algorithms are specifically designed to identify and/or consider specific patient characteristics when assisting decision-making (e.g., an algorithm intended to identify certain groups of patients susceptible to a condition or that may benefit from a particular therapy).
- HHS' proposals affecting the use of AI do not adequately consider the role of transparent communication of intended uses and related risks, and of patient consent, with respect to the appropriate use of AI tools and services by covered entities.
- Under HHS' proposal, covered entities could face liability for discriminatory outcomes realized after using an AI tool for some time, presenting a significant incentive to avoid

using AI tools altogether, which may not align with health AI-related liability distributions for other risks (e.g., patient safety).

- Machine translation tools are widely relied upon by providers, and serve as a critical tool in providing timely and efficacious care (particularly in the real-time communication context), and continue to be improved upon. HHS proposes to require a covered entity that uses machine translation to have translated materials reviewed by a qualified human translator when the underlying text is critical to the rights, benefits, or meaningful access of a limited English proficiency individual, when accuracy is essential, or when the source documents or materials contain complex, non-literal, or technical language. HHS' rationale for such a proposal lacks a sufficient evidence base of machine translation tools being blanket categorized as not fit for purpose and could effectively force any covered entity using machine translation tools to have to further provide for a human translator's review in all circumstances.
- Implementing the proposed 1557 regulations for AI will require significant efforts to build capacity within HHS to appropriately conduct fact-specific analyses of allegations of discrimination, and to work with the covered entity to achieve compliance.

To be clear, we share HHS' concerns about health AI and the impact of harmful biases, and are committed to advancing solutions to ensure that such harms are identified and mitigated. Providers, technology developers and vendors, health systems, insurers, and other stakeholders all benefit from understanding the distribution of risk and liability in building, testing, and using healthcare AI tools. We urge HHS to collaborate with all stakeholders to develop and operationalize frameworks that utilize risk-based approaches to align healthcare AI uses with consensus benchmarks for safety, efficacy, and equity. Moreover, we also urge HHS to collaborate with stakeholders to ensure the appropriate distribution and mitigation of risk and liability by supporting those in the value chain with the ability to minimize risks based on their knowledge and ability to mitigate should have appropriate incentives to do so. HHS' proposed 1557 regulatory updates for AI bias, as drafted, would derail the progress made through public-private partnerships and standardization activities, and significantly disincite covered entities' use of AI, ultimately robbing patients of the benefits of AI.

We urge OCR to build on and defer to the leading work of other federal actors with deep expertise in AI and bias risk mitigation, and to leverage these agencies' expertise to build its capacity to address AI-related concerns (e.g., training and staffing, enhanced public-private partnership activities, etc.). OCR must ensure that its AI-specific proposals in the 1557 rule build on these leading efforts, deferring to them where possible in order to avoid imposing new liability on providers with no connected public benefit. OCR's technology-neutral 1557 regulations already enable it to work with these agencies, building on their expertise, to enforce its regulations to address discriminatory outcomes that may be related to algorithms. These other federal efforts include:

- **The Food & Drug Administration (FDA)** oversees market entry for AI-based software as a medical device (SaMD) across a wide range of conditions, and provides guidance on avoiding bias in automated decision-making (e.g., the avoidance of automation bias in the context of clinical decision support⁴). FDA continues to leverage total product lifecycle oversight to further the potential that AI has to deliver safe and effective

⁴ <https://www.fda.gov/media/109618/download>.

software functionality that improves patients' quality of care.⁵ Even more recently, FDA has proposed new guidance addressing what information should be included in a Predetermined Change Control Plan that may be provided in a marketing submission for machine learning-enabled device software functions.⁶

- **The Federal Trade Commission** prohibits unfair or deceptive practices, including in the context of the sale or use of racially-biased algorithms used in healthcare.⁷
- **The Centers for Medicare and Medicaid Services (CMS)** continues to develop means to advance equity through its policies, both in the context of fee-for-service and value-based care.⁸
- **The Agency for Healthcare Research and Quality (AHRQ)** leads in reviewing, collecting, and sharing leading health and clinical evidence, including in the context of AI.⁹ AHRQ's efforts have demonstrated that reliance on a quality- and safety-focused approach to risk management results in advancement of the Quadruple Aim.
- **The United States Preventive Services Task Force (USPSTF)** has developed guidance aimed at mitigating bias harms (e.g., for addressing racism in preventative services¹⁰), which promotes a responsibility to consider direct and indirect harms in preventative care, and to manage those risks based on how likely, frequent, and severe they are.
- **The Office of the National Coordinator for Health IT** continues to collaborate with the healthcare community to develop standards for the seamless exchange of health data points for inclusion in electronic health records.¹¹ Even more recently, ONC has proposed new requirements for "decision support interventions" in the ONC Health IT Certification Program.¹²

Indeed, CHI agrees that higher-risk health AI should, per sound risk management practices, have a human being engaged in an oversight role. This approach would be consistent with the FDA's (described *infra*), and the CHI community has been forthright in supporting such measures in its health AI policy principles and subsequent detailed recommendations on good

⁵ <https://www.fda.gov/media/122535/download>; <https://www.fda.gov/news-events/press-announcements/fda-releases-artificial-intelligencemachine-learning-action-plan>.

⁶ <https://www.fda.gov/regulatory-information/search-fda-guidance-documents/marketing-submission-recommendations-predetermined-change-control-plan-artificial>.

⁷ E.g., <https://www.ftc.gov/business-guidance/blog/2021/04/aiming-truth-fairness-equity-your-companys-use-ai>.

⁸ <https://www.cms.gov/about-cms/agency-information/omh/health-equity-programs/cms-framework-for-health-equity>.

⁹ E.g., <https://effectivehealthcare.ahrq.gov/products/racial-disparities-health-healthcare/protocol>.

¹⁰ <https://www.uspreventiveservicestaskforce.org/uspstf/sites/default/files/inline-files/addressing-racism-preventive-services-technical%20brief.pdf>.

¹¹ <https://confluence.hl7.org/display/FAST/FHIR+at+Scale+Taskforce+%28FAST%29+Home>.

¹² <https://www.healthit.gov/buzz-blog/electronic-health-and-medical-records/interoperability-electronic-health-and-medical-records/oncs-new-proposed-rule-the-next-step-to-advancing-the-care-continuum-through-technology-and-interoperability>.

machine learning practices for AI medical devices as well as CHI's health AI transparency recommendations.¹³

Accordingly, CHI strongly urges HHS to withdraw its technology-specific proposal addressing covered entities' use of algorithms in its 1557 nondiscrimination rules. OCR can address such instances of bias and related discriminatory outcomes through the application of its 1557 rules across use cases that may involve new technology capabilities and means, including AI, in a technology-neutral manner to providers' activities that are in scope. OCR should explore how general 1557 regulatory language may be relied upon to address its concerns with health AI and discriminatory outcomes in a technology neutral manner. OCR should then, as a next step, undertake further consultations and evaluation of existing and ongoing research and best practices, in collaboration with other agencies and the broader healthcare community, to (1) gain understanding of the state of health AI technologies and deployments, including technical and legal realities of health technology supply chains, (2) ensure that its proposals impacting health AI and liability for discriminatory outcomes do not disincent the development and use of beneficial AI tools in healthcare, and (3) avoid misaligning liabilities for health AI-related discriminatory outcomes with the distribution of risks and liabilities related to other issues. Such a process would enable OCR to partner with our community to advance standardization and testing efforts that will mitigate AI bias harms, and contribute to the appropriate distribution and mitigation of risk and liability (*i.e.*, those in the value chain with the knowledge and ability to minimize and mitigate risks should have the appropriate incentives to do so). Such a consultation would also enable OCR and others to fully understand the affect of its AI-related proposals on covered entities' practical ability to use AI tools and services, particularly those with limited resources, and other priorities such as the need to reduce provider burnout.

If OCR nonetheless decides to retain AI-specific provisions in its new 1557 nondiscrimination rules, we strongly encourage the following steps be taken:

- ***Scope the Application of OCR's AI-Specific Provisions to Map to OCR's Goals.***
OCR's definition of AI, in its rules, should be tailored so that the rules map to the outcomes it seeks. In its present proposal, OCR would apply to "clinical algorithms" which OCR states are "tools used to guide health care decision-making and can range in form from flowcharts and clinical guidelines to complex computer algorithms, decision support interventions, and models." With no definition of "clinical algorithm" currently existing in federal regulation, application of the rules will be difficult for covered entities and could result in the application of the rules to use cases not intended by OCR. Indeed, in practice this sweeping definition could encompass nearly any technology used by a covered entity. An automated calendar tool that identifies the next available appointment for a patient or a tool that predicts when a doctor's office will need to order additional supplies are arguably within scope, even though the likelihood and severity of discrimination harm stemming from these tools is low. OCR's rule should limit the definition of covered algorithm to those that exclude human oversight and should not include technology simply used to support, inform, or facilitate decision-making. Rather than place liability on covered entities for any discriminatory algorithm, the proposed rules should be modified to impose liability where the clinical algorithm is not subject to human review or judgment to minimize the risk of over-reliance on such tools in a way that results in discrimination. If the provider's independent judgment, coupled with the use of the clinical algorithm, results in discrimination, the other provisions in Section

¹³ Connected Health Initiative, Policy Principles for Artificial Intelligence in Health, available at <https://connectedhi.com/wp-content/uploads/2022/02/Policy-Principles-for-AI.pdf>

1557 would apply and would prohibit this behavior, and the latest proposed rules' steps to reinstate and clarify what type of discrimination is unlawful would ensure that the prohibition against discrimination is meaningful.

Further, a key way to resolve definition/scope ambiguities, and to build on sister agency expertise, is to exempt uses of health AI from the scope of its rules already authorized and overseen by the FDA, as the FDA takes steps to address the biases with which OCR is concerned. OCR must recognize that application of 1557 nondiscrimination rules would present overlapping regulatory burdens in competition with expert federal partners already addressing AI bias, strongly disincenting the uptake or continued use of any AI across healthcare practices. In turn, this would adversely affect adversely America's unserved and underserved communities most acutely, undermining OCR's motivations in introducing AI-specific provisions into the 1557 regulations.

- **Clarify Reliance on an Actual Knowledge Standard.** OCR's nondiscrimination rules should align with an actual knowledge standard of liability. Providers, and others in a healthcare value chain, should only be expected to take reasonable steps to mitigate disparities in algorithms they have developed or have knowledge of. Such an approach would reflect the need for appropriate risk distribution throughout healthcare value chains, and that those in the value chain with the ability to minimize and mitigate risks based on their knowledge have appropriate incentives to do so. Therefore, we strongly encourage OCR to ensure that its rules do not assess individual liability for the performance of algorithms they did not design or know or have reason to know that the algorithm produced discriminatory results.
- **Take Measures to Enable Pro-Patient Uses of Machine Translation Tools.** OCR should revise its proposals specifically addressing machine translation to reflect the wide benefits that administrative AI functions (e.g., machine translation tools) provide today across healthcare contexts, particularly in real-time communications, and clarify that a mandate for review by a human interpreter does not apply to real-time communications (whether in-person or via video). As proposed, OCR's requirements on machine translation would result in the widespread abandonment of machine translation tools across covered entities, ultimately harming patient care, increasing healthcare costs, and adding to provider burdens. OCR's rules affecting machine translation must be predicated on a compliance analysis, weighing the net impacts of removing machine translation tools from the care continuum entirely in assessing the reasonableness of a covered entity's activities in using such machine translation tools under its proposed factors.
- **Provide a Reasonable Pathway to Compliance.** The deployer of AI-enabled technology is best positioned to understand the context and environment in which a clinical algorithm is used, and best able to manage other guardrails such as appropriate human review or judgment. In order to facilitate deployers obtaining appropriate information and assurances about the algorithms, we encourage OCR to further explore whether a HIPAA business associate-like approach could most effectively help address OCR's concerns, where entities subject to Section 1557 of the ACA would contractually require AI developers to provide assurances that the clinical algorithm does not discriminate, that appropriate steps were taken to mitigate risk of discrimination, and that the inputs of the algorithm were representative of the community in which the clinical algorithm would be used. If OCR is principally interested in creating accountability with AI developers, it could incorporate a model where entities subject to Section 1557 are

required to implement contractual terms with providers or a safe harbor from liability if terms are in place. If the HHS proposed rule were to enter into effect in its current form, covered entities are likely to create this contractual flow-down structure as a means of mitigating risk.

Further, OCR should provide a reasonable enforcement discretion period (one year) before any enforcement of AI-specific provisions would take place to enable process changes needed to comply with the rules. During such a period, we urge OCR to engage in proactive outreach and education to the healthcare community affected by its rules. And whether or not OCR does move forward with AI-specific provisions in the 1557 nondiscrimination rules, the CHI community requests OCR's assistance in understanding compliance with 1557 rules and emerging technologies, including but not limited to AI and machine learning. CHI members, and the healthcare community writ large, would benefit from compliance assistance, and we welcome the opportunity to collaborate in creating resources to address this need.

We appreciate HHS' consideration of our input on its proposals and encourage thoughtful consideration of our input. As a community, we stand ready to assist further in any way that we can.

Sincerely,

A handwritten signature in black ink, appearing to read 'Brian Scarpelli', with a stylized, cursive script.

Brian Scarpelli
Executive Director

Leanna Wade
Regulatory Policy Associate

Connected Health Initiative
1401 K St NW (Ste 501)
Washington, DC 20005