

September 22, 2023

Senator Bill Cassidy, M.D.
Ranking Member
U.S. Senate Committee on Health, Education, Labor, and Pensions
Washington, DC

RE: Exploring Congress' Framework for the Future of AI

Dear Senator Cassidy,

The Association of Clinical Research Organizations (ACRO) represents the world's leading clinical research and clinical technology organizations. Our member companies provide a wide range of specialized services across the entire spectrum of development for new drugs, biologics, and medical devices—from pre-clinical, proof of concept and first-in-human studies through post-approval, pharmacovigilance, and health data research. ACRO members manage or otherwise support a majority of all biopharmaceutical sponsored clinical investigators worldwide and advance clinical outsourcing to improve the quality, efficiency, and safety of biomedical research.

ACRO appreciates the opportunity to provide our industry's thinking and recommendations in response to your RFI on *Exploring Congress' Framework for the Future of AI*. In this response we offer some general recommendations, before addressing some of the RFI's specific questions.

General Recommendations

1. **Endpoint vs. Exploratory:** ACRO suggests that the ultimate FDA regulatory framework for AI/ML should primarily govern use cases involving AI/ML that directly impact determinants of effectiveness and safety, such as endpoint collection, patient safety improvement, and in enhancing data quality. AI/ML applications aimed at improving operational efficiency and for exploratory purposes should, in general, be excluded from the regulatory framework.
2. **Governance and Oversight:** Industry should establish strong internal governance and oversight mechanisms for monitoring the development of AI/ML models.
 - a. Create a dedicated internal governance framework for AI.
 - b. Establish dedicated teams, including an AI ethics board, responsible for implementing AI strategy.
 - c. Develop an AI handbook to integrate AI use into organizational policies.
3. **Accountability and Risk Assessments:** Effective accountability and risk assessment frameworks are necessary for controlling the deployment of AI/ML models. The Food and Drug Administration (FDA) should provide stakeholder-informed guidelines on assessing risks and implementing accountability structures. Considerations include:

- a. Develop risk assessment scales and utilize impact assessments to identify and mitigate risks.
 - b. Include AI monitoring within existing accountability structures.
 - c. Establish accountability documentation, such as impact assessments, audit functions, monitoring procedures, management plans, and AI acquisition and procurement policies. Ensure that IP and copyright related issues are assessed.
4. Human Oversight: Ensuring clear human oversight is essential for monitoring AI models and maintaining proper governance. The FDA should provide stakeholder-informed guidelines on the expected level of human oversight. Recommended areas of oversight include:
 - a. Incorporate human review and accuracy testing at each stage of the AI life cycle.
 - b. Provide oversight on the data used for training AI models.
 - c. Monitor inaccurate AI decisions and establish corrective measures managed and reviewed by humans.
 - d. Establish policies and protocols for triggering human interventions when necessary.
5. Transparency: A transparent framework is vital for governing and monitoring AI models. A uniform mandate on transparency will facilitate regulatory compliance for both the FDA and ACRO member teams. Consider the following ideas:
 - a. Articulate the organization's risk appetite, methodologies, and frameworks in an AI impact assessment document.
 - b. Draft notices, policies and procedures, disclaimer statements outlining the purpose of AI systems, underlying datasets, decision-making processes, and subject rights related to AI systems.
 - c. Ensure transparency frameworks focus on outcomes that protect patient safety and data integrity, rather than explanations of the inner workings of the systems.
6. Explainability of AI/ML Models: Research suggests that explainability and AI model performance come with certain tradeoffs.¹ ACRO believes that the totality of evidence or interpretability of the model are better indicators for good and sound model performance. Given the tradeoffs, ACRO recommends a risk-based approach to focus on outcomes and results rather than expecting detailed explainability of AI/ML tools.

¹ Forough Poursabzi-Sangdeh, Daniel G Goldstein, Jake M Hofman, Jennifer Wortman Wortman Vaughan, and Hanna Wallach. 2021. Manipulating and Measuring Model Interpretability. In Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems (CHI '21). Association for Computing Machinery, New York, NY, USA, Article 237, 1–52. <https://doi.org/10.1145/3411764.3445315>

Specific Questions

How can FDA support the use of AI to design and develop new drugs and biologics?

ACRO stresses the need for regulatory clarity—specifically in use cases that directly impact clinical trial endpoints, enhance the quality of submission data, and prioritize patient safety—in order to support the use of AI/ML in the development of new drugs and biologics. We believe that providing clear guidelines and regulations for AI/ML applications in these cases will ensure consistent standards and facilitate innovation in the life sciences industry.

However, we would like to emphasize that there are other areas where AI/ML can be utilized in clinical operations to drive efficiencies, optimize site and patient recruitment, and enhance operational processes. These use cases primarily focus on improving internal workflows, resource allocation, and operational decision-making without directly affecting regulatory submissions or patient safety. In these non-submission-related areas, we recommend that FDA refrain from regulatory oversight.

What existing standards are in place to demonstrate clinical validity when leveraging AI? What gaps exist in those standards?

To ensure the integrity and clinical validity of the use of AI/ML in clinical trials, ACRO suggests the following practices:

- **Rigorous Data Quality Assurance:** Implementing robust data quality assurance processes is crucial to ensure the integrity of AI/ML. This includes validation, cleaning, and normalization techniques to address data errors, inconsistencies, and outliers. By ensuring data accuracy, completeness, and reliability, the integrity of AI/ML models and their outputs can be maintained.
- **Validation assessments:** Conducting thorough validation assessments of AI/ML models is essential to ensure their reliability, performance, and validity for the specific context of use in clinical trials. Validation should include rigorous testing against appropriate references standards, comparison with gold-standard methodologies, and assessment of the model's generalizability and limitations. This includes implementing mechanisms to detect and mitigate biases in AI/ML models.
- **Independent Verification and Validation:** Performing independent verification and validation of AI/ML algorithms and software code helps ensure the integrity of the models. Independent experts or third-party organizations can assess the algorithms, validate the code, and verify the accuracy and reliability of the AI/ML outputs. This process enhances transparency, accountability, and confidence in the integrity of AI/ML applications.
- **Regulatory Compliance and Data Integrity Guidelines:** Incorporating regulatory compliance and adherence to relevant data integrity guidelines is essential during the development and implementation of AI/ML in clinical trials. Following established regulations and guidelines, such as those provided by regulatory authorities and industry

standards, ensures that the data used, algorithms employed, and processes followed maintain the highest standards of integrity and quality.

By implementing these practices, stakeholders in clinical trials can assure the validity and integrity of AI/ML applications, enhancing the reliability and trustworthiness of the generated insights and outcomes.

What practices are in place to mitigate bias in AI decision-making?

To address bias in AI/ML applications within the contexts of clinical trials, developers employ the following processes:

- **Bias Detection Algorithms:** Developers implement bias detection algorithms to identify potential biases in training datasets and model outputs. These algorithms analyze the data to identify patterns and discrepancies that may indicate biased outcomes or unfair treatment.
- **Fairness Assessments:** Conducting fairness assessments is crucial in evaluating the impact of AI/ML models on different demographic groups. Developers assess the fairness of model outputs across various characteristics such as race, gender, age, or socioeconomic status to identify and mitigate any identified biases. Disclosure of such assessments would help ensure that the AI/ML models do not disproportionately favor or disadvantage specific population subgroups.
- **Ongoing Model Performance Monitoring:** Developers regularly monitor the performance of AI/ML models to detect and address biases that may arise during model deployment. This monitoring involves analyzing real-world outcomes and comparing them against expected results. If biases are identified, developers can recalibrate the algorithms or introduce corrective measures to minimize their impact and ensure fair and unbiased results.
- **Collaboration with Domain Experts and Stakeholders:** Collaboration with domain experts and diverse stakeholder groups is essential to comprehensively evaluate potential biases and their implications. By engaging with experts from various disciplines and involving representatives from different demographic groups, developers can gain insights into potential biases and consider multiple perspectives in the bias identification and management process.

By implementing these processes, developers in clinical trials can enhance the identification and management of biases in AI/ML applications, promoting fairness, equity, and unbiased decision-making.

Is the current HIPAA framework equipped to safeguard patient privacy with regards to AI in clinical settings? If not, how not or how to better equip the framework?

Current HIPAA regulations provide a well-tested set of requirements for the use and disclosure of protected health information (PHI) in clinical settings and describe the obligations of covered entities (CEs) and their business associates (BAs) in the handling of

PHI. However, given the potential for unforeseen negative consequences resulting from the use of an AI application, such as an inaccurate diagnosis or an entirely incorrect medication dosage, our earlier comments about mitigation of potential biases and the importance of human oversight of AI applications (ensuring there is a “human in the loop”) are enormously important. CEs will need to be especially careful in deploying AI applications that they have not developed themselves or do not fully understand.

In terms of protecting privacy per se, AI applications must be deployed in clinical care settings only in ways that comply with HIPAA limitations on use or disclosure of PHI for any purpose other than treatment, payment and health care operations (TPO).

Conclusion

ACRO appreciates the Ranking Member’s commitment to ensuring the safe and effective use of artificial intelligence and machine learning in drug development. Thank you for considering our feedback. We are available for further discussions or clarification if needed.

Respectfully,



Sophia McLeod
Director, Government Relations
smcleod@acrohealth.org