

ConnectedHealthInitiative

October 17, 2025

Grace R. Graham
Deputy Commissioner for Policy, Legislation, and International Affairs
Food and Drug Administration
10903 New Hampshire Avenue
Silver Spring, Maryland 20993-0002

RE: Comments of the Connected Health Initiative, *Digital Health Advisory Committee; Notice of Meeting; Establishment of a Public Docket; Request for Comments-Generative Artificial Intelligence-Enabled Digital Mental Health Medical Devices* [Docket No. FDA-2025-N-2338; 90 FR 44196]

The Connected Health Initiative (CHI) writes to provide input to inform the Food and Drug Administration's (FDA) total product lifecycle considerations for Generative Artificial Intelligence (AI)-enabled devices in connection with the upcoming November 06, 2025, meeting of the Digital Health Advisory Committee (DHAC).¹

CHI is the leading multi stakeholder encouraging the use of digital health innovations and supporting an environment in which patients and consumers can see improvements in their health. We seek essential policy changes that will help all Americans benefit from an information and communications technology-enabled American healthcare system. For more information, see www.connectedhi.com.

CHI is a longtime active advocate for the increased use of new and innovative digital technologies in both the prevention and treatment of disease, and we appreciate the FDA's consistent collaboration on digital health-related technologies to responsibly streamline their pathway to the market, including through its Digital Health Advisory Committee (DHAC). AI-enabled software functions are radically improving the American healthcare system, represent the most promising avenue for improved care quality, reduced hospitalizations, avoidance of complications, and improved satisfaction, particularly for the chronically ill.

Already, AI-driven algorithmic decision tools and predictive analytics have substantial direct and indirect effects in consumer and enterprise context. 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, with drug development being no exception. As AI systems, powered by streams of data and advanced algorithms, continue to improve services and generate new business models, the fundamental transformation of

¹ *Digital Health Advisory Committee; Notice of Meeting; Establishment of a Public Docket; Request for Comments on Generative Artificial Intelligence-Enabled Digital Mental Health Medical Devices*, 90 Fed. Reg. 56174 (Sept. 12, 2025).

economies across the globe will only accelerate. Nonetheless, AI also has the potential to raise a variety of unique considerations for policymakers.

CHI appreciates the DHAC deliberating on and generating recommendations on generative AI-enabled digital mental health medical devices. These devices are already demonstrating clinical and operational value by delivering personalized, adaptive mental health support at scale, and as their algorithms and evidence base continue to mature, they are likely to further enhance care accessibility, precision, and patient engagement. We agree that, as these products further develop, new and/or novel risks may also follow, which make it critical for FDA's regulatory frameworks (including considerations for premarket evidence requirements and post-market performance monitoring) to evolve while enabling innovation that supports public health. With rising demand for mental health services and limited access to qualified providers in the United States, such new technologies will help bridge care gaps, improving both outcomes and accessibility.

CHI has worked across its diverse stakeholder community for years to proactively address Software as a Medical Device (SaMD) AI opportunities and challenges. CHI urges the DHAC to align its recommendations with the following, which are also appended to this comment letter:

- **APPENDIX 1:** CHI's *Health AI Policy Principles*, a comprehensive set of recommendations on the areas that should be addressed by policymakers examining AI's use in healthcare, and how they should be addressed (<https://connectedhi.com/wp-content/uploads/2022/02/Policy-Principles-for-AI.pdf>);
- **APPENDIX 2:** CHI's *Health AI Good Machine Learning Practices*, a recommended pathway for the FDA to ensure innovation in machine learning-enabled medical devices, including for continuously learning algorithms, while protecting patient safety: <https://connectedhi.com/wp-content/uploads/2022/04/CHIAITaskForceGMLPs.pdf>
- **APPENDIX 3:** CHI's *Advancing Transparency for Artificial Intelligence in the Healthcare Ecosystem*, a proposal on ways to increase the transparency of and trust in health AI tools, particularly for care teams and patients (<https://connectedhi.com/wp-content/uploads/2022/02/AdvancingTransparencyforArtificialIntelligenceintheHealthcareEcosystem.pdf>); and
- **APPENDIX 4:** CHI's *Health AI Roles & Interdependency Framework*, which proposes clear definitions of stakeholders across the healthcare AI value chain, from development to distribution, deployment, and end use; and suggests roles for supporting safety, ethical use, and fairness for each of these important stakeholder groups that are intended to illuminate the interdependencies between these actors, thus advancing the shared responsibility concept (<https://connectedhi.com/wp-content/uploads/2024/02/CHI-Health-AI-Roles.pdf>).

CHI appreciates the opportunity to submit its comments to the DHAC (and the FDA) and urges its thoughtful consideration of the above/attached.

Sincerely,



Brian Scarpelli
Executive Director

Chapin Gregor
Policy Counsel

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

Appendices:

Appendix A: CHI's Health AI Policy Principles

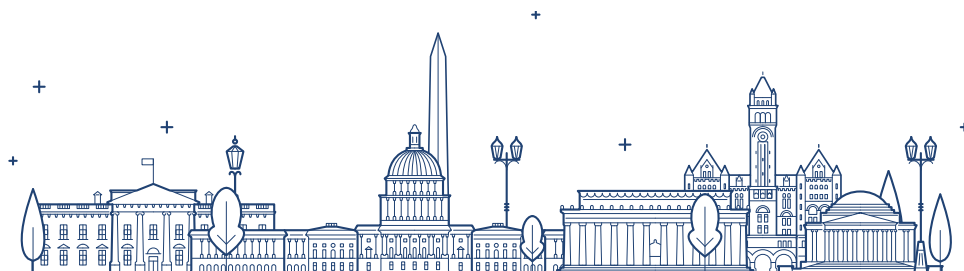
Appendix B: CHI's Health AI Good Machine Learning Practices

Appendix C: CHI's Advancing Transparency for Artificial Intelligence in the Healthcare Ecosystem

Appendix D: CHI's Health AI Roles & Interdependency Framework



Policy Principles for Artificial Intelligence in Health



Connected Health is an initiative
of ACT | The App Association

1401 K Street NW Suite 501
Washington, DC 20005

📞 202.331.2130

🌐 connectedhi.com

🐦 #connectedhealth

📘 /ConnectedHealthInitiative

Policy Principles for AI in Health

Today, there are already many examples of AI systems, powered by streams of data and advanced algorithms, improving healthcare by preventing hospitalizations, reducing complications, decreasing administrative burdens, and improving patient engagement. AI systems offer the promise to rapidly accelerate and scale such results and drive a fundamental transformation of the current disease-based system to one that supports prevention and health maintenance. Nonetheless, AI in healthcare has the potential to raise a variety of unique considerations for U.S. policymakers.

Many organizations are taking steps to proactively address adoption and integration of AI into health care and how it should be approached by clinicians, technologists, patients and consumers, policymakers, and other stakeholders, such as the Partnership for AI, Xavier Health, the American Medical Association, and the Association for the Advancement of Medical Instrumentation and BSI. Building on these important efforts, the Connected Health Initiative's (CHI) Health AI Task Force is taking the next step to address the role of AI in healthcare.

First, AI systems deployed in healthcare must advance the “quadruple aim” by improving population health; improving patient health outcomes and satisfaction; increasing value by lowering overall costs; and improving clinician and healthcare team well-being. Second, AI systems should:

- Enhance access to health care.
- Empower patients and consumers to manage and optimize their health.
- Facilitate and strengthen the relationship and communication that individuals have with their health care team.
- Reduce administrative and cognitive burdens for patients and their health care team.

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

- **National Health AI Strategy:** 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 will require strong guidance and coordination. Given the significant role of the government in the regulation, delivery, and payment of healthcare, as well as its role as steward of significant amounts of patient data, a federal healthcare AI strategy incorporating guidance on the issues below will be vital to achieving the promise that AI offers to patients and the healthcare sector. Other countries have begun to take similar steps (e.g., The UK's Initial Code of Conduct for Data Driven Care and Technology) and it is critical that U.S. policymakers collaborate with provider organizations, other civil society organizations, and private sector stakeholders to begin similar work.

- **Research:** Policy frameworks should support and facilitate research and development of AI in healthcare 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. Clinical validation and transparency research should be prioritized and involve collaboration among all affected stakeholders who must responsibly address the ethical, social, economic, and legal implications that may result from AI applications in healthcare. Further, public funding and incentives should be conditioned on promoting the medical commons in order to advance shared knowledge, access, and innovation.
- **Quality Assurance and Oversight:** Policy frameworks should utilize risk-based approaches to ensure that the use of AI in healthcare aligns with recognized standards of safety, efficacy, and equity. 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. 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 in healthcare is safe, efficacious, and equitable.
 - Ensuring algorithms, datasets, and decisions are auditable and when applied to medical care (such as screening, diagnosis, or treatment) are clinically validated and explainable.
 - AI developers should consistently utilize rigorous procedures and must be able to document their methods and results.
 - Those developing, offering, or testing healthcare AI systems should be required 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.
 - Adverse events should be timely reported to relevant oversight bodies for appropriate investigation and action.

- **Thoughtful Design:** Policy frameworks should require design of AI systems in health care that are informed by real-world workflow, human-centered design and usability principles, and end-user needs. Also, AI systems should help patients, providers, and other care team members overcome the current fragmentation and dysfunctions of the healthcare system. AI systems solutions should facilitate a transition to changes in care delivery that advance the quadruple aim. The design, development, and success of AI in healthcare should leverage collaboration and dialogue between caregivers, AI technology developers, and other healthcare stakeholders in order to have all perspectives reflected in AI solutions.
- **Access and Affordability:** Policy frameworks should ensure AI systems in health care are accessible and affordable. Significant resources may be required to scale systems in health care and policy-makers must take steps to remedy the uneven distribution of resources and access. There are varied applications of AI systems in health care such as research, health administration and operations, population health, practice delivery improvement, and direct clinical care. Payment and incentive policies must be in place to invest in building infrastructure, preparing personnel and training, as well as developing, validating, and maintaining AI system with an eye toward ensuring value. While AI systems should help transition to value-based delivery models by providing essential population health tools and providing enhanced scalability and patient support, in the interim payment policies must incentivize a pathway for the voluntary adoption and integration of AI systems into clinical practice as well as other applications under existing payment models.
- **Ethics:** Given the longstanding, deeply rooted, and well-developed body of medical and biomedical ethics, it will be critical to promote many of the existing and emerging ethical norms of the medical community for broader adherence by technologists, innovators, computer scientists, and those who use such systems. Healthcare AI will only succeed if it is used ethically to protect patients and consumers. Policy frameworks should:
 - Ensuring AI in healthcare is safe, efficacious, and equitable.
 - Ensure that healthcare 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 in healthcare, as needed.
 - Ensure consistency with international conventions on human rights.
 - Ensure that AI for health is inclusive such that AI solutions beneficial to patients are developed across socioeconomic, age, gender, geographic origin, and other groupings.
 - Reflect that AI for health tools may reveal extremely sensitive and private information about a patient and ensure that laws protect such information from being used to discriminate against patients.

- Modernized Privacy and Security Frameworks:** While the types of data items analyzed by AI and other technologies are not new, this analysis provides 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 patients. 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 health information is properly protected, while also allowing the flow of health information. This information is necessary to provide and promote high-quality healthcare and to protect the public's health and well-being. There are specific uses of data that require additional policy safeguards, i.e., genomic information. Given that one individual's DNA includes potentially identifying information about even distant relatives of that individual, a separate and more detailed approach may be necessary for genomic privacy. Further, enhanced protection from discrimination based on pre-existing conditions or genomic information may be needed for patients. 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.
- Collaboration and Interoperability:** Policy frameworks should enable eased data access and use through creating a culture of cooperation, trust, and openness among policymakers, health AI technology developers and users, and the public.
- Workforce Issues and AI in Healthcare:** The United States faces significant demands on the healthcare system and safety net programs due to an aging population and a wave of retirements among practicing care workers. And lower birth rates mean that fewer young people are entering the workforce. Successful creation and deployment of AI-enabled technologies which help care providers meet the needs of all patients will be an essential part of addressing this projected shortage of care workers. Policymakers and stakeholders will need to work together to create the appropriate balance between human care and decision-making and augmented capabilities from AI-enabled technologies and tools.
- 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. In developing and using healthcare AI solutions, these data provenance and bias issues must be addressed. 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 patients or consumers.

- **Education:** Policy frameworks should support education for the advancement of AI in healthcare, promote examples that demonstrate the success of AI in healthcare, and encourage stakeholder engagements to keep frameworks responsive to emerging opportunities and challenges.
- Patients and consumers should be educated as to the use of AI in the care they are receiving.
- Academic/medical education should include curriculum that will advance health care providers' understanding of and ability to use health AI solutions. Ongoing continuing education should also advance understanding of the safe and effective use of AI in healthcare delivery.



ConnectedHealth

Machine Learning and Medical Devices

Connecting practice to policy (*and back again*)

By Sebastian Holst with Morgan Reed and Brian Scarpelli



Machine Learning and Medical Devices

Connecting practice to policy *(and back again)*

Contents

Introduction.....	2
Effective governance is required to accelerate and amplify continued Machine Learning innovation	3
ML governance must be engineered into ML development practices and account for ML application behaviors	3
Training Data shapes ML application behavior.....	4
Source code does not predict ML application behavior	4
ML applications can continuously evolve.....	4
Effective governance of ML-enabled solutions begins with effective governance of ML software development and operations	5
Engineer effective ML governance into Medical Device software development lifecycles	5
Part 1: Trace ML-specific properties through the software development lifecycle	6
Part 2: Review Work-In-Progress: A Proposed Regulatory Framework for Modifications to AI/ML-Based Software as a Medical Device	6
Part 3: Beyond the Total Product Lifecycle	7
Tracing Machine Learning development properties through a general software development and DevOps lifecycle	9
Software Development Lifecycle Management	9
Machine Learning SDLC Requirement Summary	11
Quality Management	11
ML Software Quality Summary.....	13
Software Security and Risk Management	14
Machine Learning Security and Risk Management Summary	15
Work-In-Progress Review: A Proposed Regulatory Framework for Modifications to AI/ML-Based Software as a Medical Device.....	16
GMLP Summary	18
The Culture of Quality and Organizational Excellence	18
Culture of Quality and Organizational Excellence	20
Initial observations	21
Appendix A: Supporting organizations and underlying standards and frameworks	22
International Electrotechnical Commission (IEC)	22
International Organization for Standardization (ISO).....	22
International Medical Device Regulators Forum (IMDRF)	23
US Food and Drug Administration (FDA)	23
Appendix B: Respondent Submission Analysis	24
Proposal Questions and Feedback	24
Respondent industries and corresponding stakeholder community roles.....	25
Respondent priorities.....	25
Respondent priorities by topic	26
FDA-specific question response.....	27
Appendix C: Beyond the Total Product Lifecycle.....	28

Introduction

Today, there are already many examples of artificial intelligence (AI) systems, powered by streams of data and advanced algorithms, improving healthcare by preventing hospitalizations, reducing complications, decreasing administrative burdens, and improving patient engagement. AI systems offer the promise to further accelerate and scale such results and provide impetus to the ongoing transition from our current disease-based system to one that is centered upon prevention and health maintenance. Nonetheless, AI in healthcare also brings with it a variety of unique considerations for U.S. policymakers, particularly for medical device regulators.

Many organizations are taking steps to proactively address adoption and integration of AI into health care and how it should be approached by clinicians, technologists, patients and consumers, policymakers, and other stakeholders. Building on these important efforts, the Connected Health Initiative's (CHI) Health AI Task Force has taken the next step to address the role of AI in healthcare through the development of health AI policy principles.¹

Generally, CHI believes that AI systems deployed in healthcare must advance the “quadruple aim” by improving population health; improving patient health outcomes and satisfaction; increasing value by lowering overall costs; and improving clinician and healthcare team well-being.

In order to succeed, Health AI systems must:

- Enhance access to health care.
- Empower patients and consumers to manage and optimize their health.
- Facilitate and strengthen the relationship and communication that individuals have with their health care team.
- Reduce administrative and cognitive burdens for patients and their health care team.

In providing its health AI policy principles with various key US federal policymakers, CHI's diverse AI Task Force has identified an opportunity to expand its contribution through a projection of its health AI policy principles onto a collection of good machine learning practices (GMLPs). Through a variety of public and collaborative initiatives designed to refine and build consensus around GMLPs, the objective is to provide a baseline that the Food and Drug Administration (FDA) and other governmental and non-governmental stakeholders can leverage in their ongoing consideration of the topic. We intend for this document to serve as a next step in shaping health AI-related policy developments at the FDA, at the US federal level widely, and internationally.

CHI's AI Task Force welcomes collaboration with any interested stakeholder moving forward and appreciates consideration of this document.

¹ Connected Health Initiative *Policy Principles for Artificial Intelligence in Health*, <https://actonline.org/wp-content/uploads/Policy-Principles-for-AI.pdf>.

Effective governance is required to accelerate and amplify continued Machine Learning innovation

Machine Learning² has advanced the quality and efficiency of medical devices and promises still greater innovations at an ever-quicken pace. Machine Learning's track record coupled with sky-high expectations for the future have also spawned a proportionate demand for – and investment in – effective governance; a means of assessing Machine Learning (ML) application suitability and performance, managing associated risks, and ensuring public safety and ethical use.

This document focuses on governance with respect two primary ML system categories: continuously learning systems (CLS) that are inherently capable of learning from real-world data and are able to update themselves automatically over time while in public use and “locks down” systems that have no ability to alter their configuration once testing and certification have been completed.

Governance strives to ensure appropriate levels of transparency, reliability, safety, security, and privacy.

Effective governance delivers on these objectives without compromising utility, efficiency, or innovation.

Effective ML governance is further required to instill confidence and trust in overall quality that, in turn, will lead to increased development velocity and ever-more ambitious innovation.

ML governance must be engineered into ML development practices and account for ML application behaviors

ML software behaves differently than traditional software in large part because it is developed differently.



The fastest cars need the best brakes.
To have the confidence required to drive at the highest speeds, a driver must trust their brakes – not just for emergencies, but for every scenario and under all conditions. And, without exception, the best brakes are engineered into the car; never added on as an afterthought¹.

² CHI supports the exemplary work of numerous organizations that are addressing healthcare AI, and seeks to harmonize and build upon these efforts including reuse, wherever possible, of accepted and recognized terminology and definitions. Unless defined inline, this paper will reuse the terminology and definitions included in the December 2019-released Xavier University paper Building Explainability and Trust for AI in Healthcare. <https://www.xavierhealth.org/news3/2020/1/8>.

³ This analogy has been borrowed with gratitude from the [Open Compliance and Ethics Group](#), a non-profit think tank that promotes Principled Performance as the universal goal of every organization, team and individual.

Training Data shapes ML application behavior



Rather than explicitly define each logical sequence through source code as a traditional developer would, a ML developer transforms a generic predictive engine (an untrained machine) using a carefully curated training data set. In much the same way that a sculptor creates a mold around an original object, the ML developer creates a trained machine around a training data set. The training data set is constructed by the developer, but the training (computational analysis and resulting modifications to the untrained machine) are executed without developer intervention. The training data set has replaced source code at this stage of the development process and represents a wholly new development artifact.

How should training data sets be created, curated, and vetted?

Source code does not predict ML application behavior



There is no longer a one-to-one connection between application logic (behavior) and authored code. Depending on the training data set and the properties of the generic machine selected, the trained engine may have the ability to identify a broken bone in an X-ray, predict a heart attack, or dispense proper dosages of critical medication. Static analysis of peripheral source code or the training data set cannot predict the trained ML engine's behavior.

How can testing criteria be established if software behavior itself cannot be fully specified?

ML applications can continuously evolve



Unlike the compilation of source code into an executable program, machine training is not restricted to a single operation prior to an application's production release. If configured to do so, a trained machine that is in production (operational) can employ continuously learning systems (CLS) e.g. continue training using data consumed while in a production environment. This allows for the possibility that different copies of a single trained machine may each evolve independently from one another and from the initial trained machine.

How should new behaviors be evaluated in the field? When can this behavior even be safely deployed?

Effective governance of ML-enabled solutions begins with effective governance of ML software development and operations

The scale, complexity and distribution of ML applications has made governing each ML application instance recommendation, prediction, and action impossible.

What is possible – and practical – is to identify ML-specific risk factors stemming from the “paradigm-shifting” properties outlined above and evaluate how these have been proactively and transparently mitigated *within a broader software development lifecycle management context.*



It's not the “what”, it's the “how.”

The FDA [Food Code](#) ensures food safety and protection by focusing on broad areas of risk including the provisioning, preparation, and delivery of food.

It is not possible to evaluate each of the billions of food servings delivered every day. Governing the food supply chain and preparation “lifecycle” is the only practical means of effective governance.

How to get an A grade in ML software development

“FDA will assess the culture of quality and organizational excellence of a particular company and have reasonable assurance of the high quality of their software development, testing, and performance monitoring of their products².”



Broad Risk Categories	
Food	Machine Learning
Food from unsafe sources	Training data set deficits
Inadequate cooking	Machine training errors
Improper holding temperatures	Pipeline and distribution failures
Contaminated equipment	Operational vulnerabilities
Poor personal hygiene	Poor training and culture

Engineer effective ML governance into Medical Device software development lifecycles

There is an established practice of adapting vetted quality system management and software development lifecycle practices to support the unique priorities and requirements of the medical device industry.

The operative word here is “vetted.” Due in large part to the three paradigm-shifting properties of ML technology outlined above, general ML software quality and development practices may be, in some circumstances, less mature than the development practices currently in place. The potential immaturity of some ML quality and risk management practices suggests that something more than “adapting” generally accepted practices will be necessary.

Given the accrued history and expertise of today’s healthcare software developers – and SaMD developers in particular – this community has a material contribution to make in advancing – not merely adapting – mainstream development best practices.

⁴ Proposed Regulatory Framework for Modifications to Artificial Intelligence/Machine Learning (AI/ML)-Based Software as a Medical Device (SaMD)

Part 1: Trace ML-specific properties through the software development lifecycle

The first task is to consider where traditional software development and quality management practices are most likely to require ML-specific accommodations prior to suggesting follow-on medical-device-specific adjustments.

The approach taken here is to trace ML-specific properties through the software development lifecycle. In much the same way that a contrast MRI employs a dye to highlight specific and difficult to detect conditions, this paper traces ML-properties across three interwoven software development axes with a special sensitivity to healthcare's overriding priorities, e.g. safety, transparency, and accuracy. The three development axes are:

1. Software manufacturing (the general principles of how whatever is developed is constructed, delivered, and maintained),
2. Software quality management (how suitability of purpose is defined and assessed for what is manufactured), and
3. Software security and risk management (frameworks and practices for identifying, assessing, and mitigating risks stemming from missed manufacturing or quality management requirements).



A Contrast MRI

A contrast MRI uses the injection of a contrast dye to better highlight certain conditions that might otherwise go undetected.

Part 2: Work-In-Progress Review: A Proposed Regulatory Framework for Modifications to AI/ML-Based Software as a Medical Device

In April of 2019, The FDA published an ambitious work that incorporated ML-centric principles into existing software development practices⁵, [Proposed Regulatory Framework for Modifications to AI/ML-Based Software as a Medical Device \(SaMD\) - Discussion Paper and Request for Feedback](#).

The stated goal was to advance a framework that would allow the FDA's regulatory oversight to embrace the iterative improvement power of machine learning for Software as Medical Device while assuring that patient safety is maintained.

Safety assurance is achieved through a multi-pronged approach that includes recommendations that ensure ongoing ML algorithm changes are:

- Implemented according to pre-specified performance objectives,
- Follow defined algorithm change protocols,
- Utilize a validation process that is committed to improving the performance, safety, and effectiveness of AI/ML software, and



Proposed Regulatory Framework for Modifications to Artificial Intelligence/Machine Learning (AI/ML)-Based Software as a Medical Device (SaMD)

Discussion Paper and Request for Feedback



⁵ The authors acknowledge their debt to the International Medical Device Regulators Forum (IMDRF) for their work on SaMD (which, itself, relies upon prior IEC and ISO standards and frameworks) while recognizing the need for a “new, total product lifecycle (TPLC) regulatory approach that facilitates a rapid cycle of product improvement and allows these devices to continually improve while providing effective safeguards.”

- Include real-world monitoring of performance.

These recommendations are rolled into an updated Total Product Lifecycle (TPLC) regulatory framework with the ultimate aim of promoting a mechanism for manufacturers to be “continually vigilant in maintaining the safety and effectiveness of their SaMD,” supporting “both FDA and manufacturers in providing increased benefits to patients and providers.”

As with The Food Code, the FDA would assess the culture of quality and organizational excellence of a particular company in order to establish “reasonable assurance” of the high quality of their software development, testing, and performance monitoring of their products.

Given that general-purpose software development practices are themselves undergoing a material ML-driven evolution,

- Are there any underlying assumptions regarding quality and audit that merit closer review?
- What assurances can be built-in to ensure that those changes will be appropriately reflected in the central regulatory notions of “a culture of quality and excellence” and “reasonable assurance?”

Part 3: Beyond the Total Product Lifecycle⁶

Are there untapped approaches to embrace ML’s most dynamic and opaque (but potentially powerful) properties? Are there longer-term opportunities to reimagine certification and pre-certification roles and workflows to further leverage AI/ML innovations?

Perhaps the most radical ML property from a regulatory perspective is the potential for algorithms to evolve after release and distribution. This capability is what is referred to as continuously learning systems.

Currently, this is only a theoretical concern as there is a blanket prohibition of this scenario across every existing and proposed TPLC regulatory framework.

Might there come a time when this prohibition will be perceived as imposing an undue constraint on innovation? Is there a scenario – perhaps in a robotics context – where allowing an initial set of SaMD instances to evolve wholly independently from one another will be identified as an absolute requirement? How would today’s notions of manufacturing lifecycle and quality need to adapt?

The FDA, Machine Learning & SaMD

The FDA’s has already begun the complex task of reimagining regulatory oversight to best embrace the power of machine learning while continuing to assure patient safety.



Only “frozen algorithms” need apply (for now)

As with a graduating class of identically trained physicians whose skills mature independently over time, it is possible for an initial set of ML SaMD instances to evolve wholly independently from one another after distribution.

Might there come a time when the prohibition of real-time, continuous learning is perceived as an undue constraint on innovation?

⁶ See Appendix C: Beyond the Total Product Lifecycle

Machine Learning is not the only transformative computing force. Cloud services, mobile 5G, and blockchain are among a growing list of revolutionary technological domains that are enabling entirely new ways of working, collaborating, and communicating.

Are there near-term organizational or technological opportunities that can help to prioritize near-term ML regulatory, governance and compliance requirements while also better positioning stakeholders across the healthcare and technology spectrum to capitalize on what may appear at first to be ML's most radical properties?

Tracing Machine Learning development properties through a general software development and DevOps lifecycle

Healthcare software governance combines policies and controls to:

- Ensure public safety
- Mitigate risks stemming from
 - Unintended consequences
 - Poor execution
 - Adversarial exploitation
- Encourage innovation in applications as well as the specialized development and testing tools required to produce those applications.

In what ways might ML development properties challenge foundational assumptions underlying traditional development lifecycle management practices?

Software Development Lifecycle Management

Software Development Lifecycle (SDLC) Management and DevOps tooling and practices normalize and automate software manufacturing processes while helping to ensure that safety, transparency, and privacy requirements are met.

In order for Machine Learning to complete its transition from paradigm-shifting innovation to a mainstream technology, SDLC management must also meet any additional requirements stemming from ML data-driven machine training development practices, e.g. Machine Learning Software Development Lifecycle Management (MLDLC).

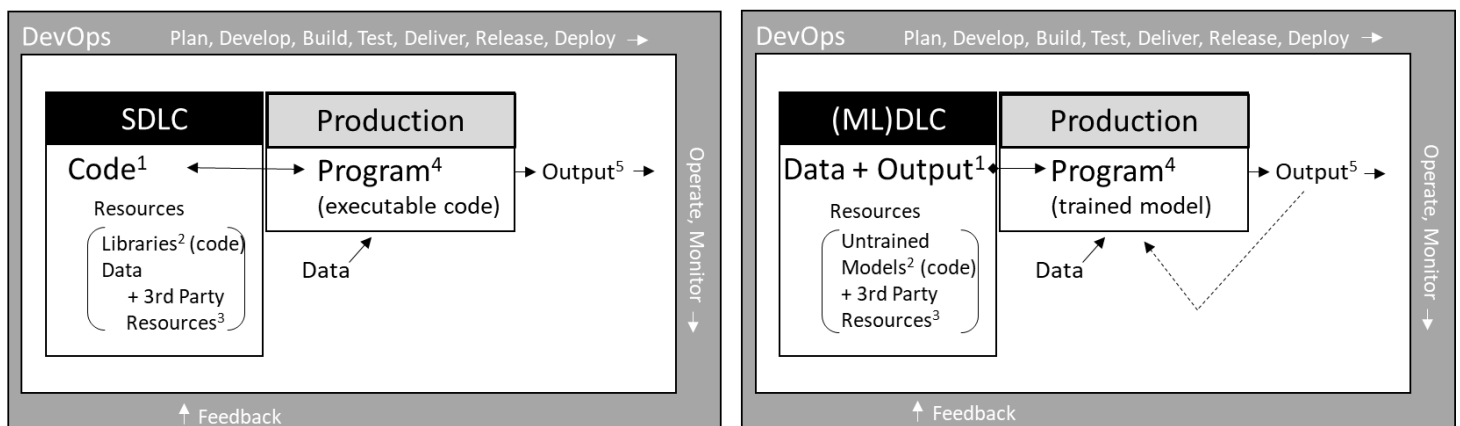


Figure 1: Traditional SDLC versus Machine Learning MLDLC wrapping in a DevOps iterative pipeline.

Figure 1 illustrates the elements of, and relationships between, a traditional Software Development Lifecycle and a Machine Learning Development Lifecycle operating within a well-formed DevOps pipeline.

The Figure 1 notes are described in the following table.

ML Key	Topic	Note
  	1 Code vs Data + Output	Code sits at the center of a traditional SDLC and, consequently, is subject to rigorous quality, audit, and sourcing controls. Given that Data + Output supplants Code in a MLDLC, it follows that <i>an equivalent – but not identical – collection of controls are needed to ensure that effective quality, audit and sourcing remain in place.</i>
	2 Libraries vs Untrained Models	A traditional SDLC has built-in support for managing reusable code, typically in the form of libraries, to speed and simplify development, improve quality and auditability, and to help ensure consistency over time and across development teams. In much the same fashion, MLDLC will draw from a collection of reusable untrained models ⁷ . These models are code-based and are often organized as a traditional library, but given their heightened impact on development outcomes, <i>a corresponding increase in Untrained Model governance may also be justified.</i>
	3 3 rd Party Resources	Today's applications increasingly rely upon 3 rd party managed services, libraries, and software components. SDLC tools (Integrated Development Environments or IDE's) as well as software and service distribution channels have been extended to better support this rapidly evolving software supply chain. Supply chain risk management has also evolved to ensure appropriate visibility and accountability as the sourcing of code and services become increasingly distributed and diverse. <i>IDE's and IT security and risk management frameworks must evolve in-kind to keep pace with the consequences of including 3rd party Data + Output and/or Untrained Models into the modern software supply chain.</i>
	4 Production Programs	The traditional SDLC deliverable is an executable program. The MLDLC deliverable is a trained model. Due to ML statistical techniques, it is typically not possible – or nearly impossible – to trace exactly why a trained model behaves as it does. The absence of a decision tree in an ML program renders traditional SDLC code reviews, debugging, and general monitoring techniques obsolete. <i>ML programs may require compensating mechanisms to ensure comparable degrees of transparency, reliability and auditability.</i>
	5 Output	Both traditional SDLC and DevOps best practices include a feedback loop that can be used to generate new requirements or improve existing features. This kind of continuous feedback fuels future program iterations and is subjected to the complete SDLC beginning with requirements through coding, test, etc. However, there are some branches of Machine Learning, specifically Continuously Learning where feedback is delivered directly into the current Production ML Program. These classes of Machine Learning bypass traditional SDLC inspection and approval steps and may result in unplanned and, potentially, unexpected behaviors. <i>Owners and regulators of sensitive and high-risk applications that must include human inspection may need to consider a blanket prohibition of these subcategories of Machine Learning until new norms about acceptable risk and transparency can be established. At a minimum, a greater understanding of the limitations and side-effects of deployed machine learning algorithms will be required by auditors and regulators.</i>

Table 1: MLDLC requirements stress traditional SDLC practices.

⁷ ML programs also include “traditional reusable code” as well.

Machine Learning SDLC Requirement Summary

Tracing ML properties through high level SDLC stages suggested several potential new or modified requirements including:

1. The transition from code-driven to data-driven development will require corresponding practices and controls to meet quality, audit, and sourcing requirements.
2. Reusable Untrained Models are a special class of reusable code that, given their heightened impact on development outcomes, require a proportionate increase in governance.
3. Security and risk management must evolve in-step to keep pace with the implications of including 3rd party Data + Output and/or Untrained Models into the modern software supply chain.
4. Production ML programs may require novel monitoring and debugging mechanisms to ensure acceptable transparency, reliability, and auditability
5. Owners and regulators of sensitive and high-risk applications may need to consider blanket prohibitions of CLS Machine Learning models unless and until revised notions of transparency and predictability are established.
6. Integrated Development Environments (IDE's) and associated tooling will need to be extended to better scale and automate all phases of the new MLDLC.

Quality Management

While SDLC management measures and manages software manufacturing, distribution, and consumption, Software Quality is the field of study and practice that describes, measures, and manages the desirability (suitability) of the software itself.

Production Software Quality is, in large part, built upon Software Program Quality (the executable) that is, in turn, built upon the underlying Code Quality.

The shift to trained models away from code suggests a requirement to supplement existing code-centric quality practices and metrics.

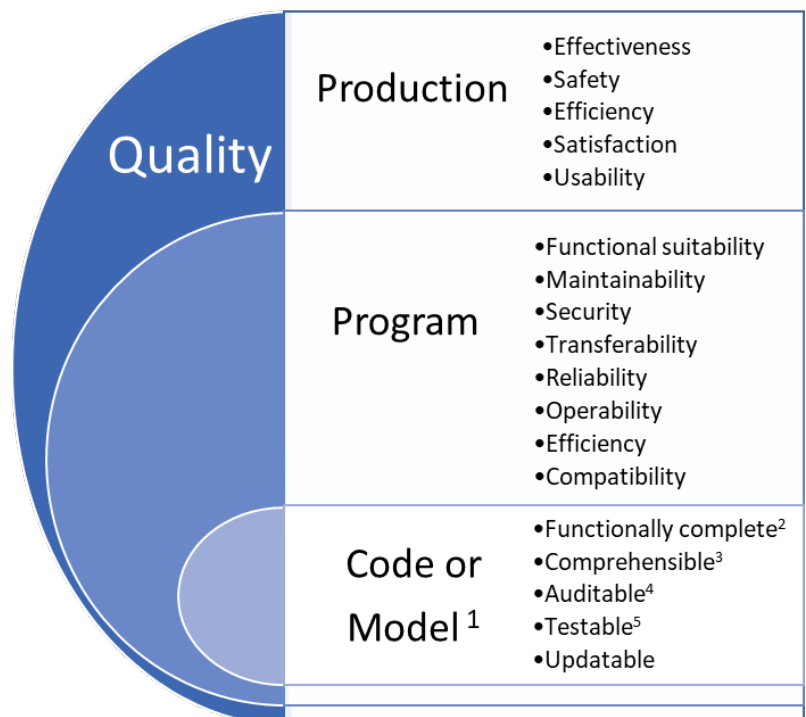


Figure 2: Quality is managed throughout the development lifecycle.

Figure 2 illustrates the elements of, and relationships between, common quality metrics divided into three segments: underlying code (or trained model), the resulting program, and the performance or suitability of that program.

Figure 2 notes are described in the following table:

ML Key	Quality Topic	Note
	1 Code vs Trained Model	Code sits at the center of a traditional Software Quality Practice with well-defined subcategories including functional completeness, comprehensibility, auditability, testability, and updatability. To preserve overall Quality, <i>ML development must develop equivalent – but not identical – methods of measuring and establishing acceptable quality metrics and tolerances.</i>
	Trained Model vs Code	
	2 Functionally Complete	Code can be statically analyzed, monitored for “coverage”, and otherwise exercised to generate a mapping of input data and environmental states to expected outcomes. ML models are trained and tested through the processing of carefully curated data sets – there is no code that can be parsed and traced. Poorly formed datasets generate unexpected and potentially unpredictable, behaviors and/or incorrect weighting of outcome predictions. Common examples of training data set gaps include: - Insufficient data volume - Lopsided data distribution across activities and outcomes - Missing activities and/or outcomes - Impossible activities or outcomes Poor data sets can result in the compromise multiple functional subcategories including: - Suitability: will the software behave appropriately for all users? - Accuracy: are functions implemented correctly? The models themselves may meet the highest quality standards, but the resulting trained model may fail to meet those standards. - Compliance: is the software in compliance with the necessary laws and guidelines? Transparency and predictability are required with virtually every regulatory and/or compliance obligation. <i>Development must have reliable means of detecting and, as needed, remediating gaps and other data set irregularities prior to ML model training.</i>
	3 Comprehensible	Every ML model includes intrinsic limitations. Understanding the stated purpose and objectives of a ML application and the hosting platform and implementation language will not be sufficient to assess the suitability of either training data or the selected ML models. In order to meaningfully “comprehend” the expected behavior of a trained model, <i>a reviewer must have specialized data science expertise and be knowledgeable in the strengths and limitations of the applied model(s) and the data staging/cleansing/sampling techniques.</i>
	4 Auditable	Tracing, reverse-engineering, and predicting how a model will behave given a specific set of inputs is difficult and, in practical terms, often impossible. This is especially true with extremely complex systems with many thousands of variables; the most common examples include image recognition, robotics, and natural language processing. <i>A consensus on acceptable alternatives to traditional event logging in code-based applications are needed to provide a comparable degree of assurance.</i> Untrained models are often provided by open source communities or platform providers. <i>A common format for sourcing the precise model and version with a record</i>



		<i>of know Quality issues would help to predict Quality issues that may arise in the final trained model.</i>
5	Testable	Exception detection, defect definition, and related KPI's (including testing cost) must be established to effectively model the severity and cost of ML application defects specifically related to under-performance. Output measurement must also be standardized, utilizing what developers measure for their own data models including terminology and their own interpretation of medical information. <i>This industry-specific formulation results in a harmonization of terminology across regulators and stakeholders that will improve quality management.</i>

Table 2: Trained Models drive expansion of code-centric Software Quality practices.

ML Software Quality Summary

Tracing ML properties across basic Quality System segments suggested several additional new or modified requirements including:

1. ML Software must meet the same quality standards as code-based software. As such, there must be equivalent methods of measuring and establishing acceptable ML-centric quality metrics and tolerances to offset inapplicable code-centric controls.
2. ML-centric controls must cover both the special data sets used for training and testing ML models as well as the trained ML models themselves.
3. Reviewers, testers, and auditors will require additional specialized data science expertise including a working knowledge of the strengths and limitations of deployed model(s), the implications of their parameters as well as any data staging/cleansing/sampling techniques that are applied.
4. The sourcing of untrained models is a potential supply chain gap – in much the same way that a revised compiler can introduce quality issues in established source code. A common format for sourcing a precise model and version with a record of known quality issues would likely help to predict Quality issues that may arise in a final trained model.
5. Quality Systems must also incorporate updated and harmonized health care specific terminology, data collection, and measurement practices to ensure the availability of relevant baseline healthcare quality metrics and standards.
6. The establishment of exception detection, defect definition, and related KPI's (including testing cost estimation) are needed to effectively model the severity and cost of ML application defects specifically related to ML under-performance.

Software Security and Risk Management

Effective risk and security management begins with identifying and prioritizing material threats and works to establish effective controls that reduce risk to acceptable levels. For application risk and security management, recommended practices typically include:

- Detailed Abuse Cases⁸ that are used to
 - Develop a business/technical specific Threat Model⁹ that in turn is used to assess risks stemming from
 - Each application's Attack Surface¹⁰, e.g. the application's entry and exit points.

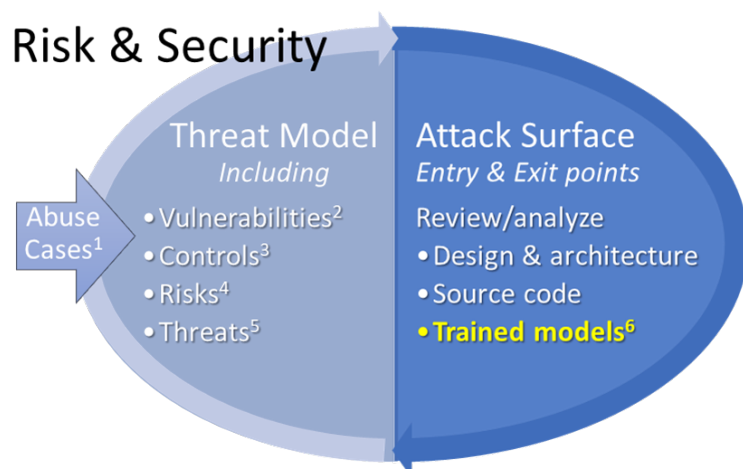


Figure 3: Risk and Security Modeling

These interrelated components evolve with production usage and feedback generating additional Abuse Cases that in turn update the Threat Model resulting in further refinements to the application's Attack Surface and underlying controls.

Software Security and Risk Management practices must also expand to meet new requirements stemming from Machine Learning development practices, technology, and use cases. Figure 3 notes are described in the following table.

ML Key	Risk & Security	Note
1 Training Data Set ML Learning	1 Abuse Cases	<i>The current paucity of established ML Abuse Cases is likely to lead to an incomplete view of potential threats and undermine threat modeling activities and the subsequent control priorities that follow.</i>
1 Training Data Set ML Learning	2 Vulnerabilities	ML systems novel use of training data to create production behaviors have spawned an equally novel set of novel vulnerabilities including: • Data poisoning (injecting training data designed to cause errors) • Adversarial input (data crafted to be misclassified by targeted models) • Exploitation of errors in autonomous system goals <i>The set of known ML-specific vulnerabilities is almost certainly incomplete as are the range of potential exploits.</i>
1 Training Data Set ML Learning	3 Controls	<i>There is a further deficit in established Preventative and Detective Controls to mitigate the risks stemming from ML-inspired vulnerability attacks.</i>
	4 Risks	Effective risk assessments are dependent upon accurate probability estimates. Risk calculations typically combine:

⁸ OWASP Abuse Case Cheat Sheet

https://github.com/OWASP/CheatSheetSeries/blob/master/cheatsheets/Abuse_Case_Cheat_Sheet.md

⁹ OWASP Application Threat Modeling

[https://www.owasp.org/index.php/Application_Threat_Modeling#1. What are we building.3F](https://www.owasp.org/index.php/Application_Threat_Modeling#1._What_are_we_building.3F)

¹⁰ OWASP Attack Surface Cheat Sheet

https://github.com/OWASP/CheatSheetSeries/blob/master/cheatsheets/Attack_Surface_Analysis_Cheat_Sheet.md



		• The probability of an incident occurring (an exploit of a vulnerability) • The probability of that incident causing harm and • The degree of harm that comes with each occurrence <i>The rapidly evolving use of ML across industries and use cases significantly complicate ML risk assessment calculations making risk mitigation investment decisions more difficult to calibrate.</i>
	5 Threats	In addition to the exploitation of unique ML vulnerabilities, the weaponization of ML in the hands of bad actors must also be considered. Examples include: • Automation of social-engineering attacks and the dissemination of political misinformation leveraging improved profiling, messaging and deep fake image and audio generation. • Anonymization and scaling of physical assaults using autonomous drones and other vehicles • Highly efficient and distributed cyber-attacks leveraging specialized ML models. • Expansion of potential attackers as democratization of all of the above removes human domain expertise as a requirement. <i>ML expands the variety of potential threats, improves the efficiency of existing threats, and expands the number of potential attackers.</i>
	6 Trained models	<i>ML training and test data sets represent additional attack surface opportunities to be included in current Attack Surface mapping practices.</i>

Table 3: Machine Learning impact on established Application Risk and Security practices

Machine Learning Security and Risk Management Summary

Tracing ML properties through security and risk management categories highlight some measure of risk from all three ML property categories listed above.

- | | |
|---|---|
| 1. The short history of successful ML exploits constrains Threat Modeling practices. | making risk mitigation investment decisions more difficult to calibrate. |
| 2. The inventory of ML-specific vulnerabilities is incomplete as are the understanding of potential exploits. | 5. ML training and test data sets represent additional attack surface opportunities to be included in current Attack Surface mapping practices. |
| 3. There is a further deficit in established Preventative and Detective Controls to mitigate the risks stemming from ML-inspired vulnerability attacks. | 6. ML has a multiplicative effect on Risk and Security management by expanding the variety of potential threats, improving the efficiency of existing threat tactics, and expanding the number of potential attackers |
| 4. The rapidly evolving use of ML across industries and use cases significantly complicate ML risk assessment calculations | |

Work-In-Progress Review: A Proposed Regulatory Framework for Modifications to AI/ML-Based Software as a Medical Device

[A Proposed Regulatory Framework for Modifications to AI/ML-Based Software as a Medical Device \(SaMD\) - Discussion Paper and Request for Feedback](#) was published with the stated goal of advancing a framework to allow the FDA’s regulatory oversight to embrace the iterative improvement power of machine learning for Software as Medical Device while assuring that patient safety is maintained.

The proposed Total Product Lifecycle (TPLC) regulatory framework is designed to ensure ongoing ML algorithm changes are:

- Implemented according to pre-specified performance objectives,
- Follow defined algorithm change protocols,
- Utilize a validation process that is committed to improving the performance, safety, and effectiveness of AI/ML software, and
- Include real-world monitoring of performance.

In order to manage the scale and scope of this ambitious effort and to avoid the necessity of auditing every development milestone of every software component, the FDA proposes assessing the culture of quality and organizational excellence of a particular company in order to establish “reasonable assurance” of the high quality of their software development, testing, and performance monitoring of their products.

As outlined in the prior section, much of the underlying general-purpose software development standards, frameworks, and practices¹¹ are themselves actively undergoing their own ML-driven evolution. This section drills into the updated Total Product Lifecycle Regulatory approach and the associated “Culture of Quality and Organizational Excellence” to identify:

- Underlying assumptions regarding Software Development Lifecycle Management, Quality or Risk that may merit closer review, and
- Mechanisms to ensure evolving assumptions are appropriately reflected in the central notions of “a culture of quality and excellence” and “reasonable assurance.”

In order to “balance the benefits and risks, and provide access to safe and effective AI/ML-based SaMD,” the revised TPLC seeks to establish clear expectations on quality systems and good ML practices (GMLP) as outlined in the following illustration.

¹¹ See Appendix A: Supporting organizations and underlying standards and frameworks.

GMLP Summary

Evaluating GMLP in the context of the ongoing evolution of ML-centered development quality, SDLC, and risk management, the following issues may merit deeper investigation:

1. Heavy reliance on standards that have historically been defined by methodical and deliberate revision policies may not be able to keep pace with rapidly changing development practices and exacerbate rather than mitigate quality risk stemming from ML's data-driven versus code-driven properties.
2. Without a sufficient body of verified ML development patterns have been documented, it may be difficult to establish a durable definition of "reasonable" and "effective."
3. The long-standing requirement that all copies of a given device or software instance can only be updated but cannot independently evolve prohibits a subset of dynamic and continuously learning applications.
4. Incident management and platform monitoring systems will likely need to expand incident categories and severity ratings to account for unique classes of exceptions unique to ML services.

The Culture of Quality and Organizational Excellence

The Culture of Quality and Organizational Excellence is itself comprised of three management principles:

1. Leadership that sets the organizational tone,
2. Lifecycle Support Processes that wrap and operationalize the actual development, and at its core,
3. Deployment, and maintenance activities associated with actual SaMD development.

As noted in Table 4, note 1 above, software lifecycle standards, such as IEC 62304, are code centric and will likely need to be extended or adapted to the unique lifecycle requirements associated with training ML algorithmic models.



FDA SaMD QMS Principles

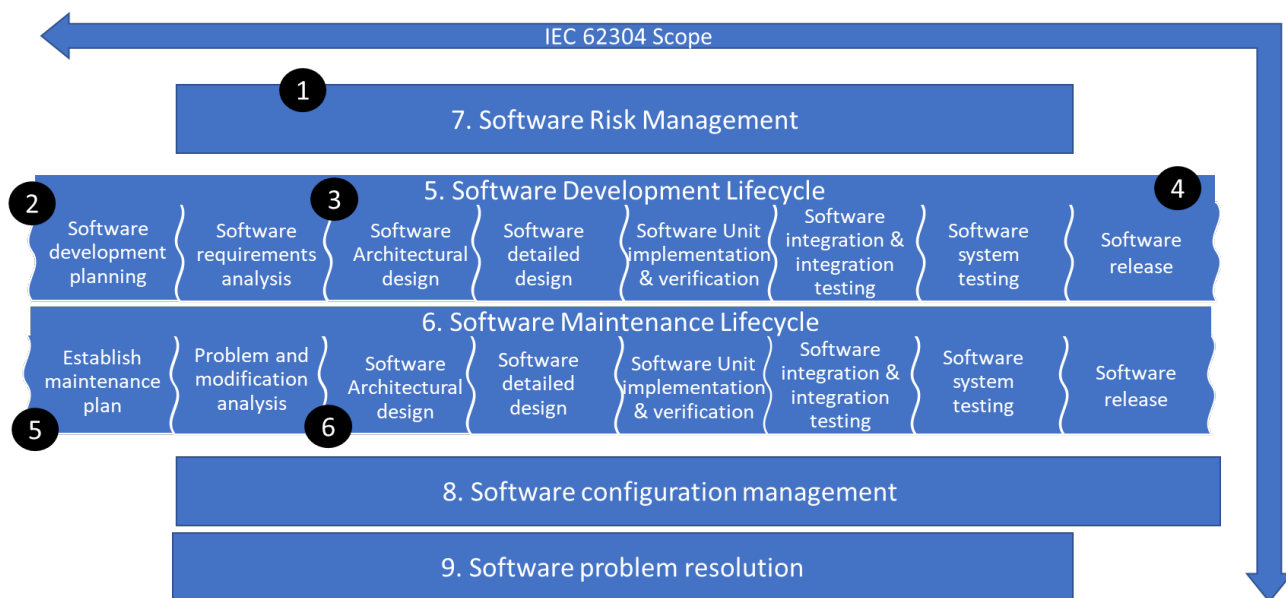


Figure 5: ML development impact on IEC 62304 development lifecycle processes.

Figure 5 notes are described in the following table.



Note	ML development impact on IEC 62394 development lifecycle processes.
1	Software Risk Management: <i>How will risks associated with training data sets be mitigated?</i>
2	Software development planning: <i>How will Software Of Unknown Providence (SOUP) be extended to accommodate 3rd party algorithms and external training data?</i>
3	Software requirements analysis: <i>How will issues relating to bias and transparency be incorporated?</i>
4	Software release: <i>Given the requirements above, how can FDA Premarket Safety Assurance requirements be effectively be met?</i>
5	Maintenance plan: <i>Defining, measuring, and documenting the degree of change within an SaMD will require significant coordination and consensus.</i>
6	Problem and modification analysis: <i>Documenting root causes and effectivity of modifications stemming from data set deficiencies will require new (or enhanced) concepts, tooling and terminology.</i>

Table 5: ML development considerations within IEC 62304: Medical device software lifecycle processes.

Culture of Quality and Organizational Excellence

Evaluating working definition of the Culture of Quality and Organizational Excellence in the context of the ongoing evolution of ML-centered development quality, SDLC, and risk management, the following issues may merit deeper investigation:

1. To satisfy an external auditor/examiner, Organizations will need to be able to tap into a sufficiently large body of recognized ML controls able to substantially meet their requirements.
2. Suppliers of third party and embedded software, also referred to as Software Of Unknown Provenance (SOUP), must be able to satisfy corresponding requirements for transparency, safety, security, and privacy.
3. Individuals will need the ability to know if/how their data may be used to develop and/or train machines or algorithms. The opportunity to participate in data collection for these purposes must be on an opt-in basis.^{12 13}
4. A consensus must be reached on the definition and measurement of a wholly new quality criteria related to behavior, e.g. bias and human-readable decision-making transparency.
5. New (or enhanced) concepts, tooling and terminology will likely be required across a broad spectrum of operations management capabilities to properly capture the impact of dataset deficiencies including:
 - Chance control documentation including risks assessment,
 - Root cause analysis, and
 - Modification effectiveness.

¹² Connected Health Initiative *Policy Principles for Artificial Intelligence in Health*, <https://actonline.org/wp-content/uploads/Policy-Principles-for-AI.pdf>.

¹³ American Medical Association's privacy principles <https://www.ama-assn.org/system/files/2020-05/privacy-principles.pdf>.

Initial observations

There is wide agreement that existing regulations need revision to accommodate the unique (and potentially disruptive) properties of Machine Learning technologies and development processes.

100% of Proposed Regulatory Framework responses endorsed the requirement to update existing medical device regulatory obligations to accommodate Machine Learning¹⁴.

The FDA, responding to this need, has proposed a regulatory framework to manage what is likely to be one of the most challenging aspects of regulating ML-driven “Software as Medical Devices,” modifications that may, or may not, require a review and recertification – a potentially time-consuming and expensive process.

One of the distinguishing properties of the Machine Learning approach is the capacity for programs to alter behavior over time without requiring additional coding or software updates. This kind of unsupervised learning challenges conventional development, quality, and risk practices and policies.

The FDA proposal built off existing regulations, frameworks, and definitions, extended some where needed, and added wholly new constructs when it was determined to be unavoidable.

Initial feedback to the proposed framework reinforced the importance of leveraging existing standards and framework – perhaps to an even greater extent than the initial proposal envisioned.

There is significantly more work that needs to be done refining and harmonizing definitions, completing core processes and performance metrics, as well as educating the vast community of stakeholders.

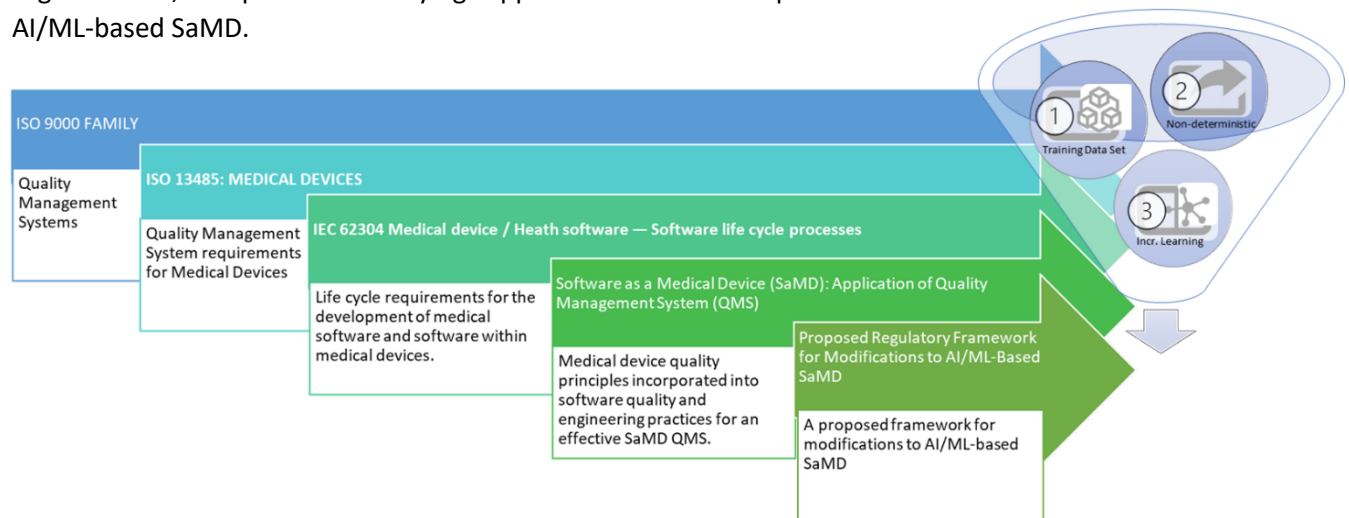
Tracing ML-specific development and technical properties from Innovator practices through relevant tooling, development frameworks, and standards promises to ultimately shorten and simplify the work required to effectively and efficiently “protecting the public health by ensuring Software as Medical Device safety, efficacy, and security.”

This can be most effectively accomplished through a sustained collaboration with, and communication across, the stakeholder ecosystem (innovators, platform providers, supranational standards bodies, government regulators, etc.).

¹⁴ See Appendix B: Respondent Submission Analysis

Appendix A: Supporting organizations and underlying standards and frameworks

There is an established practice of adapting vetted quality system management and software development lifecycle practices to support the unique priorities and requirements of the medical device industry. The following list includes frameworks and documents, as well as the associated governing organizations, that provide underlying support for the FDA's Proposed Framework for Modifications to AI/ML-based SaMD.



International Electrotechnical Commission (IEC)

The IEC prepares and publishes International Standards for all electrical, electronic and related technologies.

[IEC 62304:2006/AMD 1:2015](#) Medical device software life cycle processes is a standard which specifies life cycle requirements for the development of medical software and software within medical devices.

[International Organization for Standardization](#) (ISO)

ISO is an independent, non-governmental international organization with a membership of 164 national standards bodies. ISO – in conjunction with the IEC – has identified the need to develop standards for AI that “can benefit all societies.” Established in 2017, this is the charter of the ISO/IEC Joint Technology Committee (JTC) 1 / Subcommittee (SC) 42 for artificial intelligence (SC 42).

SC 42's scope includes basic terminology and definitions, risk management, bias and trustworthiness in AI systems, robustness of neural networks, machine-learning systems and an overview of ethical and societal concerns. SC 42 has already published three Big Data standards with 13 projects currently under development. Five of these are highlighted below.

[ISO/IEC JTC 1/SC 42](#): Artificial Intelligence

AI/ML ISO standards under development from ISO/IEC JTC 1/SC 42 include:	
ISO/IEC 23053	Framework for Artificial Intelligence (AI) Systems Using Machine Learning (ML)
ISO/IEC 24027	Bias in AI systems and AI aided decision making
ISO/IEC 38507	Governance implications of the use of artificial intelligence by organizations
ISO/IEC 23894	Artificial Intelligence — Risk Management
ISO/IEC TR 24368	Artificial Intelligence (AI) — Overview of ethical and societal concerns

[International Medical Device Regulators Forum](#) (IMDRF)

The IMDRF is a voluntary group of medical device regulators from around the world who have come together to form the Global Harmonization Task Force on Medical Devices (GHTF) whose mission is to “accelerate international medical device regulatory harmonization and convergence.” Their relevant works to date are highlighted here.

IMDRF publications include:	
IMDRF/SaMD WG/N10	SaMD: Key Definitions
IMDRF/SaMD WG/N12	SaMD: Possible Framework for Risk Categorization & Corresponding Considerations
IMDRF/SaMD WG/N23	SaMD: Application of Quality Management System
IMDRF/SaMD WG/N41	SaMD: Clinical Evaluation

[US Food and Drug Administration](#) (FDA)

The FDA is responsible for protecting the public health by ensuring the safety, efficacy, and security of drugs, biological products, *and medical devices*. In addition to the [Proposed Regulatory Framework for Modifications to AI/ML-Based Software as a Medical Device \(SaMD\) - Discussion Paper and Request for Feedback](#), the FDA is also active in contributing to, endorsing, and re-publishing many of the IMDRF publications listed above. At this time, the FDA has not made ML-specific modifications to Medical Device regulatory obligations (see [21 CFR Parts 803 through 861](#)).

Appendix B: Respondent Submission Analysis

Proposal Questions and Feedback

While there were no constraints placed on the kinds of feedback or questions that could be submitted, the FDA included questions that covered the most important (or perhaps controversial) elements of the proposed TPLC framework.

Questions included in Proposed Regulatory Framework were divided into subtopics.

- How complete is the classification of AI/ML SaMD modifications and will they be effective and helpful?
- Is the GMLP complete? How can the FDA help manufactures incorporate new requirements into their existing QMS systems and practices?
- All feedback to the definitions and implementation details surrounding SPS and ACP. These are entirely new elements to the proposed certification process.
- How can the process of premarket review (review prior to an initial SaMD launch) be better defined and managed?
- How can “real-world” data be captured, analyzed, secured, and weighted throughout this entire process?
- What should the ACP include and how can it be consistently and effectively assessed across manufacturers and SaMDs?

These questions bring to the fore just how potentially disruptive Machine Learning may be in the short-term – and why it is in everyone’s interest to shorten the ML transition into the mainstream.

That being the case, why did 64% of respondents fail to answer even one of the FDA’s questions?

64% of the public responses did not directly reference a single question included in the Framework Proposal.

Questions included in Proposed Regulatory Framework

The types of AI/ML-SaMD modifications (Key: **AI/ML SaMD**)

1. Do these categories of AI/ML-SaMD modifications align with the modifications that would typically be encountered in software development that could require premarket submission?
2. What additional categories, if any, of AI/ML-SaMD modifications should be considered in this proposed approach?
3. Would the proposed framework for addressing modifications and modification types assist the development AI/ML software?

Good Machine Learning Practices (Key: **GMLP**)

1. What additional considerations exist for GMLP?
2. How can FDA support development of GMLP?
3. How do manufacturers and software developers incorporate GMLP in their organization?

SPS and ACP (Key: **SPS/ACT**)

1. What are the appropriate elements for the SPS?
2. What are the appropriate elements for the ACP to support the SPS?
3. What potential formats do you suggest for appropriately describing a SPS and an ACP in the premarket review submission or application?

Premarket review (Key: **PreMarket**)

1. How should FDA handle changes outside of the “agreed upon SPS and ACP”?
2. What additional mechanisms could achieve a “focused review” of an SPS and ACP?
3. What content should be included in a “focused review”?

The transparency and real-world performance monitoring (Key: **Transp & Monitoring**)

1. In what ways can a manufacturer demonstrate transparency about AI/ML-SaMD algorithm updates, performance improvements, or labeling changes, to name a few?
2. What role can real-world evidence play in supporting transparency for AI/ML-SaMD?
3. What additional mechanisms exist for real-world performance monitoring of AI/ML-SaMD?
4. What additional mechanisms might be needed for real-world performance monitoring of AI/ML-SaMD?

ACP Scope: (Key: **ACP**)

1. Are there additional components for inclusion in the ACP that should be specified?
2. What additional level of detail would you add for the described components of an ACP?

The following analysis is based upon the public responses to The Proposed Regulatory Framework for Modifications to AI/ML-Based Software as a Medical Device (SaMD) - Discussion Paper and Request for Feedback.

Looking at the respondents' own questions and/or their interest (and/or lack of interest) in the FDA's questions offers insight into how stakeholders outside of the FDA perceive these issues and which of these may be perceived as more (or less) important or controversial.

Respondent industries and corresponding stakeholder community roles

Respondent submissions are available for review on the FDA website¹⁵. Figure B1 maps the self-identified Industry Categories of 127 respondents to generic Stakeholder Community roles¹⁶.

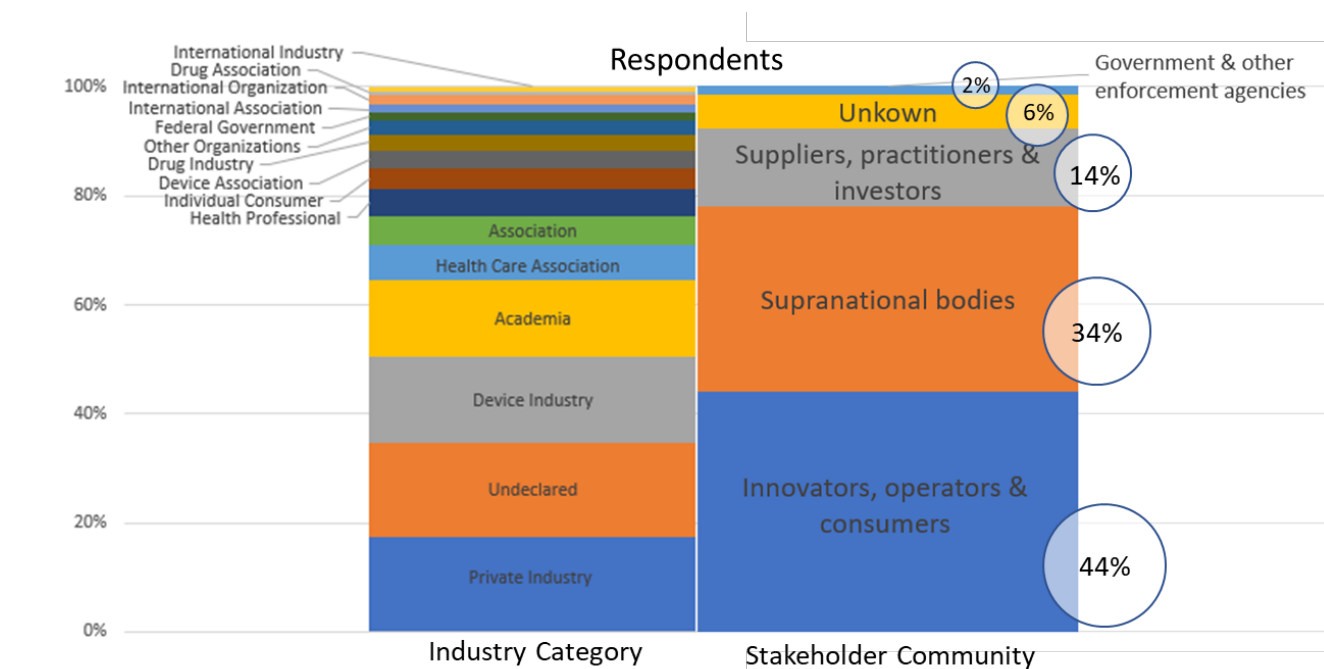


Figure B1: Respondent Industry Categories and Stakeholder Roles

Perhaps it is not surprising to learn that the primary stakeholders have the loudest voice (at least by sheer volume), but, given the importance of vendor-neutral, independent “Supranational bodies” in shaping regulations, should they?

Respondent priorities

The questions embedded inside the FDA’s regulatory framework proposal are calibrated to address the FDA’s priorities, but are those priorities and their relative weighting shared? Figure B2 illustrates the percentage of responses that included specific topics. These topics are grouped into “framework-specific” (that are unique to the proposed regulatory framework) and “mainstream activities” (that are general issues already described relating to the mainstreaming of any disruptive technology).

¹⁵ <https://www.regulations.gov/document?D=FDA-2019-N-1185-0001>

¹⁶ When included, the respondent’s organization was also used to map into the Stakeholder Community role.

64% of respondents did not answer any of the 18 questions included in the proposal. Closer inspection of respondents’ comments suggests a difference in emphasis and, perhaps, priority.

Respondents that did answer FDA-specific questions:

1. Were much more likely to comment on the ML SaMD modification categories, the recertification criteria and process, and the description of the TPLC.
2. Consistently raised issues across the mainstream activities of Quality, Risk, Ecosystem (collaboration across roles) and Frameworks (reconciliation with other frameworks).
3. Respondents that did not answer the FDA-specific questions were significantly more likely to focus on software Quality and Risk issues.
4. Regardless of whether the FDA-specific questions were addressed, there was a general concern around the definition and treatment of “Locked” models.

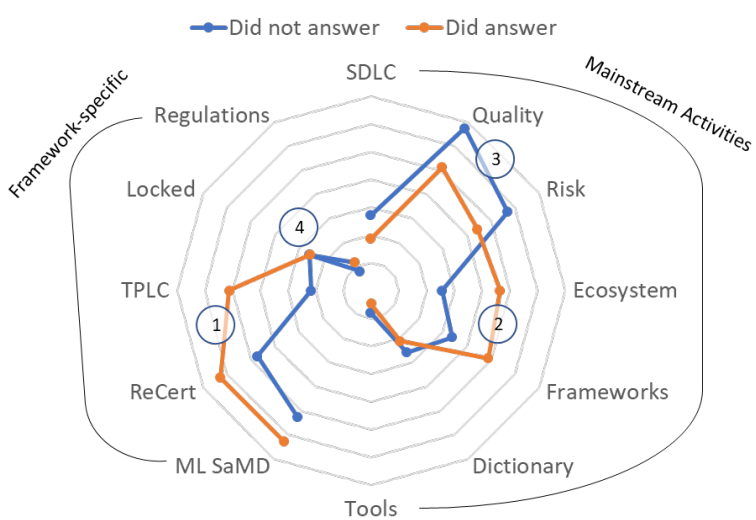


Figure B2: Topic interest of respondents

Respondent priorities by topic

Does a respondent’s stakeholder role as innovator or standards body (versus regulatory agency or consumer) also influence their priorities? If yes, should the dominance of one stakeholder role over all others be factored-in or weighted when considering responses?

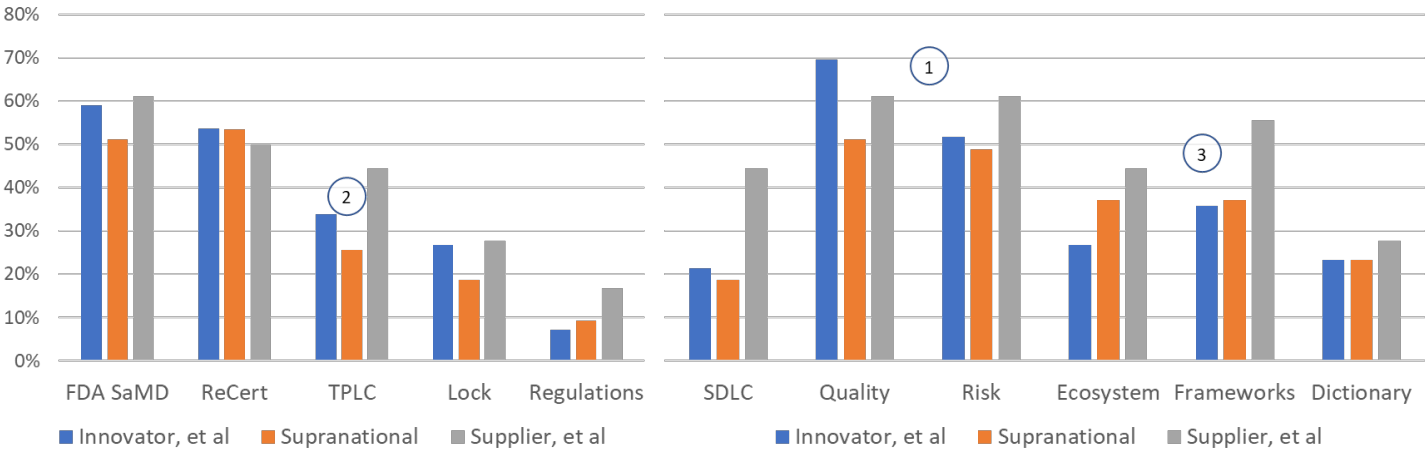


Figure B3: percentage of responses across topics by Ecosystem Stakeholder role.

Figure B3 maps the percentage of topics included in responses by Stakeholder role (only three roles had enough responses to be statistically meaningful).

1. Quality, Risk, FDA SaMD modifications and recertification processes received the greatest attention.
2. Generally, Innovators, consumers, practitioners and suppliers responded more consistently with one another as compared to Supranational organization responses.

3. Taken as a group, comments relating to Ecosystem (cross roll collaboration), Frameworks (cross framework reconciliation), and Dictionary (defining common terms and definitions across domains) were a strong, consistent area of concern.

FDA-specific question response

While only 36% of respondents addressed the embedded 18 questions directly, those responses were extensive and, obviously, important to assess.

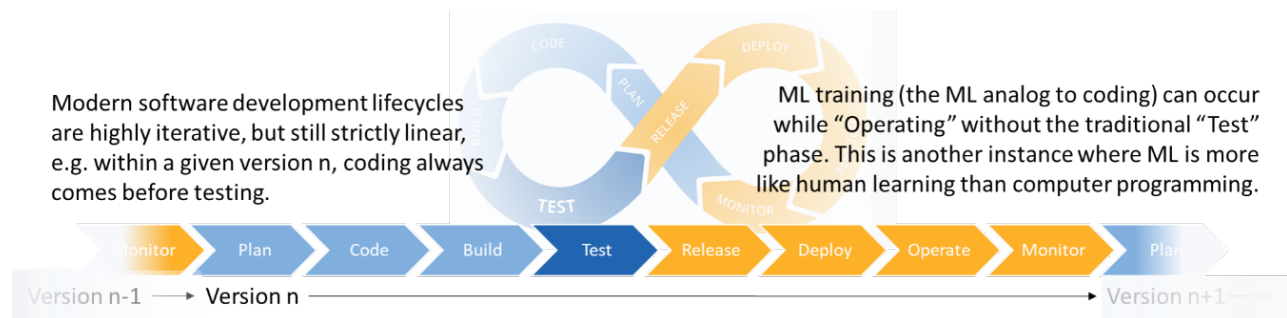


Figure B4: Count of responses that included commentary for each FDA-embedded question. The questions are segmented by topic. All Respondents are shown alongside the three highest reporting Ecosystem Stakeholder roles.

1. Respondents gave the greatest amount of attention to the questions relating to Good Machine Learning Practices.
2. Relative to the other subtopics, Algorithm Change Protocol received substantially less attention from Innovators, et al than any other subtopic. This gap was not evident in either of the other two Stakeholder roles.
3. The high innovator response volume depressed the relative importance of the ACP subtopic. Given the close relationship between Supranational Organizations and Government Regulators already discussed and the consensus around the importance of framework and regulatory consistency, should the (apparent) lack of interest from Innovators be discounted?

Appendix C: Beyond the Total Product Lifecycle

Software development lifecycle management, like virtually all modern Product Lifecycle Management, is a highly iterative process, but within any given version, the lifecycle stages are executed in a strictly linear sequence. As an example, within a given version n, coding, building, and testing must always



precede deployment and production operation.

When configured to do so, continuously learning algorithms can breach the strict sequencing imposed by development lifecycle methodology. Not surprisingly, the FDA’s proposed AI/ML TPLC includes a prohibition of this kind of evolutionary behavior in real-time and in production. This is a sound policy as there is no precedent to contradict this position to be found in the underlying standards and frameworks.

Yet, while there is no *underlying* precedent, might there be a precedent to be found in an *adjacent health care domain*?



Who’s Who and What Do They Do?

To assure patient safety, every healthcare worker must, on a reoccurring basis, be credentialed by an array of professional, State and Federal agencies.

Expensive and time consuming: Credentialing costs the U.S. healthcare system billions of dollars per year and it is time consuming. Credentialing one physician takes, on average, 100 days; a time period where that physician cannot practice.

Thanks to encrypted digital ledgers, mobile technology, and cloud services, this seemingly intractable bureaucratic nightmare is being reimaged and rebuilt as a high-speed, on-demand service able to support existing regulatory and statutory obligations at scale – improving patient safety and increasing healthcare professional availability.

If this technology can be trusted to credential hundreds of thousands of mobile healthcare professionals – what would it take to credential and authenticate millions of continuously learning medical devices?

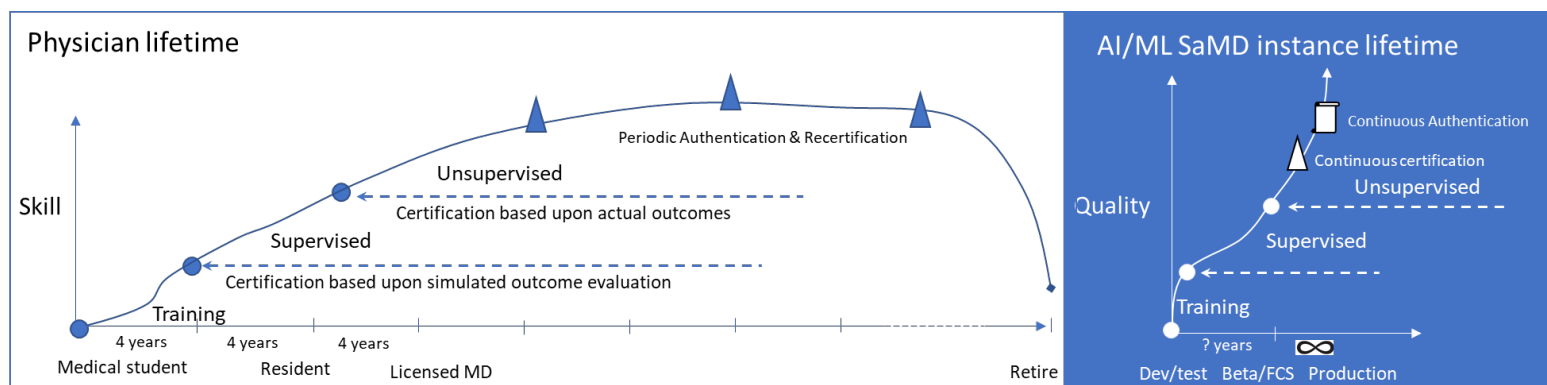


Figure C1: modeling an individual SaMD instance Quality as an independent healthcare worker's Skill.

The training, testing, and certification of a physician is not unlike the (ML)DLC or the FDA's GMLP for an AI/ML SaMD. The two only truly diverge after "certification." A physician is expected to continue to learn and improve – often in ways that are distinct from other physicians who were part of the same graduating class, a.k.a. the same release.

While there are governing bodies and controls in place to monitor the maturation of each individual physician – and to remove their privileges when needed – AI/ML SaMDs cannot be monitored individually today. As such, to assure patient safety, individual SaMD instance continued growth cannot be permitted.

Could a similar technology cocktail of encrypted digital ledgers (blockchain), mobile, and cloud technologies scale to reliably authenticate and then certify each individual medical device instance?

The first question that needs to be asked and answered is what innovation or benefits will be lost if continuous learning in production cannot be deployed. If there is no compelling use case, subsequent issues around monitoring and regulating their safety are moot.

What is evident is that, in order to remain relevant and support innovation, every interested party must remain open to reimagining the traditional roles and relationships between innovators, regulators, patients, service providers, et al alongside the coming waves of ML discoveries and breakthroughs.

Advancing Transparency for Artificial Intelligence in the Healthcare Ecosystem

OCTOBER 2021

Connected Health is an initiative of ACT | The App Association

1401 K Street NW Suite 501, Washington, DC 20005
202.331.2130 | connectedhi.com

 #connectedhealth

 /ConnectedHealthInitiative

Executive Summary

Today, the most well-known FDA-approved applications of artificial intelligence and machine learning (AI/ML) technology in healthcare are diagnostic tools that help clinicians read and interpret images to predict, detect, and monitor a number of diseases, including diabetic retinopathy and lung cancer. In the future, the use of AI/ML technology in both operational and clinical settings promises to enable a more proactive approach to healthcare that promotes investments in preventative care that can result in fewer hospitalizations, fewer doctor visits, and fewer treatments. Across use cases, AI/ML technology is helping, and must increasingly help, the healthcare industry move away from a reactive disease treatment approach to a population health management approach that lowers costs and improves care.

The immense potential of AI/ML technology in healthcare may never be fully achieved, however, unless AI/ML technologies first earn the trust of healthcare professionals and patients. The cornerstone of building trust in AI/ML technologies is to enhance transparency – providing sufficient and appropriate information about the AI/ML, including its intended use, development, performance, and, when available, logic. The more understandable the decision-making process is for each individual technology, the more confidence there will be in AI/ML use in the healthcare system.

The recommendations in this Connected Health Initiative (CHI) AI Task Force report, informed by a public roundtable CHI held to address AI/ML transparency and extensive consultations with stakeholders from across the digital health ecosystem, represent a holistic approach to creating and maintaining the trust of both healthcare professionals and patients. The Task Force set out the foundational steps AI/ML tool developers must take to build transparency into their products, but it also outlines the important roles that clinicians, healthcare providers, regulators, academic medical institutions, and accrediting organizations must play.

The medical and technology communities have a shared responsibility to provide caregivers and patients (as well as other stakeholders) with an assurance of quality through truthful representations clearly indicating the AI/ML's intended uses and risks that would be reasonably understood by those intended and expected to use the AI/ML. Uptake will depend on the buy-in of clinicians who first develop trust in AI/ML software as a medical device (SaMD) through use and experience, establishing confidence as it is adopted into practice. Once adopted, clinicians can then work with their patients to explain their use of SaMD AI/ML and inspire the same trust and confidence from the patients in the output of the SaMD AI. Each step in this chain requires buy-in and support from policymakers (both within and outside of government).

The foundation of any successful use of AI/ML technologies in healthcare depends on the trust of healthcare professionals and patients, and we believe these recommendations present a clear path toward earning that trust.



About the Connected Health Initiative

CHI is the leading multistakeholder policy and legal advocacy effort driven by a consensus of stakeholders from across the connected health ecosystem. We aim to realize an environment where Americans can improve their health through policies that allow for connected health technologies to enhance health outcomes and reduce costs. Having members who are developers and users of connected health technologies across a wide range of use cases, CHI serves as an active advocate before Congress, numerous U.S. federal agencies, and state legislatures and agencies. We seek to advance responsible pro-digital health policies and laws in areas including reimbursement and payment, privacy and security, effectiveness, and quality assurance, U.S. Food and Drug Administration (FDA) regulation of digital health, health data interoperability, and the rising role of artificial intelligence and machine learning (AI/ML) in care delivery.

In 2019, CHI formed a Task Force focused on policy challenges and opportunities related to the use of AI/ML in healthcare. CHI's AI/ML Task Force already developed a set of health AI/ML policy principles addressing how policy frameworks should adopt the role of AI/ML in healthcare.¹ A cornerstone of these principles is the idea of requiring those developing, offering, or testing healthcare AI/ML systems to provide truthful representations clearly indicating the intended use and risks that would be reasonably understood by those intended and expected to use the AI/ML solution. Such steps will provide much-needed quality assurances to caregivers and patients (as well as other stakeholders) and assist in resolving data issues that arise when an algorithm is fed bad data that can skew its learning and introduce bias. CHI's AI Task Force later developed detailed Good Machine Learning Practices for FDA-regulated AI,² which reflect and elaborate on this priority. The recommendations in this paper build on those deliverables.

Numerous CHI Steering Committee members and other key stakeholders from throughout the healthcare value chain participate in this Task Force and share a commitment to realizing the value of AI/ML in healthcare while protecting patient safety and advancing the quadruple aim. The recommendations in this paper find basis in an evaluation by the Task Force of the healthcare ecosystem's implementation of AI/ML to date, challenges and opportunities reflected by federal policymakers, and the existing and emerging issues created by AI's deployment. This report is also informed by a CHI public roundtable held in April 2021 on how to improve AI/ML transparency for caregivers and patients based on their needs and concerns, during which a wide range of stakeholders contributed to a discussion exploring novel approaches to transparency of AI/ML taken today.

For more information, please visit www.connectedhi.com.

¹ <https://actonline.org/wp-content/uploads/Policy-Principles-for-AI.pdf>.

² <https://bit.ly/3B6nslm>.

Artificial Intelligence's Role in a Successful Healthcare Ecosystem Requires Transparency

Responsible implementation of AI/ML in healthcare leads to improved medical outcomes and overall increased cost savings

Today, there are many important operational and clinical AI/ML solutions in use and many more in development.³ Some of the most well-known applications of AI/ML in healthcare that have received market clearance from the FDA are diagnostic tools that help clinicians read and interpret images. For example, AI/ML image analysis software can assist clinicians in predicting, detecting, and monitoring a number of diseases, including diabetic retinopathy, lung cancer, prostate cancer, and skin cancer. Such AI/ML uses are generally intended to be used to assist human clinicians in providing more efficient and accurate results, rather than autonomously diagnosing disease.

Separately, research projects within and outside of clinical settings continue to further explore AI's potential to revolutionize healthcare. For example, an AI/ML system developed by researchers at Northwestern University's Feinberg School of Medicine correctly identifies small lung cancer tumors nearly 95 percent of the time, while radiologists undertaking the same task unassisted are correct only 65 percent of the time.⁴ Researchers at Carnegie Mellon developed a miniature mobile robot called HeartLander that uses machine learning algorithms to make treating ventricular fibrillation (VF)—a deadly type of cardiac arrhythmia that requires cardioversion and then, if the patient survives, surgical removal of faulty heart tissue—far safer and less invasive.⁵

As a recent research paper discussing challenges related to deployment of AI/ML technologies into the clinical setting stated, “the success of a deep learning model does not rest solely on its accuracy.”⁶ The researchers noted that clinician “experiences with the system, and the socio-environmental factors that impacted system performance” must be evaluated and addressed for these systems to function in the clinical setting with the accuracy rates illustrated in the lab setting.⁷ Clearly, if the challenges of integrating AI/ML tools into clinical workflow can be overcome, AI/ML can support clinicians in a wide range of other areas. Its potential to reshape the healthcare landscape is profound, especially in the improvements it can bring to any process within healthcare operation and delivery.

Medical devices and systems that use AI/ML also represent a real opportunity to drive down healthcare costs for consumers, practitioners, and healthcare businesses alike. It is estimated that AI/ML applications can cut annual U.S. healthcare costs by \$150 billion by 2026.⁸ Most of these cost reductions stem from changing the healthcare model from a reactive to a proactive approach, focusing on health management rather than disease treatment. This focus on using AI/ML as an investment in

3 The FDA now publicly lists AI/ML medical devices cleared for marketing in United States, and includes their intended uses. See <https://www.fda.gov/medical-devices/software-medical-device-samd/artificial-intelligence-and-machine-learning-aiml-enabled-medical-devices>.

4 <https://www.nature.com/articles/d41586-020-03157-9>

5 <https://onlinelibrary.wiley.com/doi/10.1002/rcs.2297>

6 Emma Beede et al, A Human-Centered Evaluation of a Deep Learning System Deployed in Clinics for the Detection of Diabetic Retinopathy, CHI Conference on Human Factors in Computing Systems (April 2020) available at <https://dl.acm.org/doi/fullHtml/10.1145/3313831.3376718>.

7 *Id.*

8 <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC7325854/>.

preventative care can result in fewer hospitalizations, fewer doctor visits, fewer treatments, and thus fewer side effects. AI-based technology will have an important role in helping people stay healthy via remote monitoring technologies and coaching and will ensure earlier diagnosis, tailored treatments, and more efficient follow-ups.⁹

For example, AI/ML image analysis technologies can reduce medical expenses in several ways. For one, AI/ML systems can be very helpful in augmenting a clinician's analysis and treatment decisions more quickly. AI/ML technologies enable clinicians to provide the same, accurate service in a fraction of the time, increasing the volume of patients without increasing time spent treating them.¹⁰ Second, a patient whose disease is diagnosed early will pay less to treat or cure the disease than one who catches it later. The longer a disease goes undiagnosed, the more damage it causes and more resources it takes to treat, assuming it remains treatable at all. Wearable technologies that use AI, such as remote monitoring technologies, increase access to healthcare and increase engagement in treatment plans by, for example, analyzing user health data in real time and notifying wearers or their healthcare providers (or both) of potential health issues.

By introducing new, accurate, and timely data streams for human clinicians' review, AI/ML medical tools and systems that use wearable technologies can enable practitioners to come up with care and treatment options without having to see a patient in person as much, reducing administrative and in-office visit resource expenditures, and, during outbreaks of communicable diseases, at lower risk of infection to both provider and patient. The use of such technologies will also enhance patient engagement in their own care plans. This same concept also applies to laboratory technologies that use AI/ML systems, where the work hours currently required for repetitive and routine tasks could see drastic reductions, significantly cutting labor costs.¹¹

Increased efficiency, precision, and affordability are just some of the benefits that AI/ML can offer the healthcare community and those they serve, but realizing these benefits will depend on the buy-in of the provider and patient communities as well as support for responsible deployments from policymakers. CHI's AI/ML Task Force released detailed policy principles,¹² as well as proposed good machine learning practices for AI/ML meeting the definition of a medical device,¹³ to address these challenges. Notably, CHI's AI/ML Task Force has acknowledged that without its processes being understandable by humans and transparency (providing sufficient and appropriate information about the AI/ML, including its intended use, development, performance, and, when available, logic), particularly for patients and caregivers, AI/ML cannot most effectively improve healthcare. Namely, those developing, offering, or testing healthcare AI/ML systems must provide truthful and understandable representations regarding intended use and risks that would be reasonably understood by those intended, as well as expected, to use the AI/ML software as a medical device (SaMD) solution.

9 *Id.*

10 See McPhail et al, Stage at diagnosis and early mortality from cancer in England (Br J Cancer 2015), doi: [10.1038/bjc.2015.49](https://doi.org/10.1038/bjc.2015.49).

11 Rong, et al, "Artificial Intelligence in Healthcare: Review and Prediction Case Studies," Engineering, doi: [10.1016/j.eng.2019.08.015](https://doi.org/10.1016/j.eng.2019.08.015) at Sec. 2.2.

12 <https://actonline.org/wp-content/uploads/Policy-Principles-for-AI.pdf>.

13 <https://bit.ly/3B6nslm>.

How Can Transparency into Healthcare AI/ML Solutions be Advanced?

While evidence of healthcare AI's potential for widespread benefit continues to build, that potential can never be realized without healthcare professionals and patients understanding and trusting AI/ML solutions. The more transparent the decision-making process is for each individual technology, the more confidence there will be in AI/ML use in the healthcare system.¹⁴ Transparency for healthcare AI's intended uses must happen at several levels, disseminating tailored messaging to specific audiences that require insights into the AI/ML solution to make informed decisions. Building the trust that must be a foundation for the responsible deployment of AI/ML is a shared responsibility amongst developers, providers, and regulators.

Providing transparency into health AI/ML must start with the developers of the AI/ML tools. Then, uptake of AI/ML will need to be built on the buy-in of clinicians who first develop trust in AI/ML SaMD through use and experience, establishing confidence as it is adopted into practice. Once adopted, the provider can then work with his or her patients to explain their use of SaMD AI/ML and inspire the same trust and confidence by the patient in the output of the SaMD AI. Each step in this chain requires buy-in and support from policymakers (both within and outside of government).

The CHI AI/ML Task Force's recommendations for enhancing transparency for health AI/ML include:

Developers of AI/ML SaMD should:

- Prioritize making healthcare AI/ML solutions reasonably safe, efficacious, and equitable from the earliest stages of design, considering the perspectives of both patients and providers, leveraging and where necessary tweaking medical AI/ML guidelines on research and ethics,¹⁵ leading standards,¹⁶ and other resources as appropriate.
- Employ algorithms that produce repeatable results and, when feasible, are auditable, and make decisions that, when applied to medical care (such as screening, diagnosis, or treatment), are clinically validated and where possible understandable using rigorous procedures with documented methods and results, fostering efficacy through continuous monitoring.
- Rigorously identify, disclose, and mitigate biases in datasets used to train algorithms.
- Utilize risk-scaled privacy protection mechanisms for patients' data to account for the fact that the analysis by health AI/ML tools provides greater potential utility of those data items to other individuals, entities, and machines, providing many new uses for, and ways to analyze, the collected data, as well as correspondingly stronger incentives for malefactors to attempt to obtain access unlawfully. Specific uses of data that require additional safeguards (such as genomic

¹⁴ <https://www.bsigroup.com/globalassets/localfiles/en-gb/about-bsi/nsb/innovation/mhra-ai-paper-2019.pdf>

¹⁵ E.g., World Health Organization, 'Ethics & Governance of Artificial Intelligence for Health' (2021), available at <https://www.who.int/publications/i/item/9789240029200>.

¹⁶ E.g., Consumer Technology Association, 'The Use of Artificial Intelligence in Health Care: Trustworthiness (ANSI/CTA-2090)' (2021), available at <https://shop.cta.tech/collections/standards/products/the-use-of-artificial-intelligence-in-healthcare-trustworthiness-cta-2090>.

information) may necessitate a tailored approach or enhanced protections from discrimination (e.g., pre-existing conditions or genomic information may be needed for patients).

- Comply with all applicable legal and regulatory requirements.
- Develop a tailored communications and engagement plan that gives patients and providers representative of the AI/ML tool's user group a reasonably justifiable level of confidence in healthcare AI's efficacy. Such communications should enable these patients and providers to visualize the AI, and to receive direct and clear information about how their health data are being collected and used (while also avoiding information overload) and how biases in data that exacerbate disparities in healthcare are being mitigated. Reflecting that the division of labor between the developers of AI-enabled tools and the clinician or patient is critical, clearly explain intended uses, including whether a tool might include the restriction that it is not for diagnostic use or for informational purposes only, as well as risks.

Providers should:

- Develop their own risk-based and tailored communications and engagement plan that enables them to explain to patients the development of the AI/ML application, its maintainance, its performance, and how it aligns with the latest best practices and regulatory requirements to improve patient safety using easily understood and standardized formats. Providers should also acknowledge that “best practices” are dynamic and prone to obsolescence.
- Offer further detail for patients in additional resources that explain the clinical testing of AI/ML applications and the confirmation of the results by clinical experts.

The Food and Drug Administration (FDA) should:

- Leverage its successful approach to authorizing medical device AI¹⁷ that has already safely brought health AI/ML innovations to patients and providers to develop a comprehensive regulatory approach to AI/ML that meets the definition of a medical device. The FDA can accomplish this by, for example, progressing its Software Precertification Pilot¹⁸ to a full program available to all developers of SaMD AI, FDA can also update its rules and processes to realize its envisioned total product lifecycle (TPLC) regulatory approach, facilitating a potentially rapid cycle of product improvement and allowing these devices to continually improve while providing effective safeguards. This new approach should leverage CHI's Good Machine Learning Practices to address both locked and continuously learning AI.
- Evolve its requirements on reporting type and frequency so that such requirements can be adapted and scaled based on relevant factors such as risk, extent, and magnitude of

17 Software as a Medical Device (SaMD): Clinical Evaluation:

<https://www.fda.gov/downloads/medicaldevices/deviceregulationandguidance/guidancedocuments/ucm524904.pdf>; Deciding When to Submit a 510(k) for a Software Change to an Existing Device: <https://www.fda.gov/downloads/medicaldevices/deviceregulationandguidance/guidancedocuments/ucm514737.pdf>.

18 Pre-Cert Program Version 1.0 Working Model:

<https://www.fda.gov/downloads/MedicalDevices/DigitalHealth/DigitalHealthPreCertProgram/UCM629276.pdf>.

modifications, and the demonstrated reliability of the AI (e.g., quality control plans for updates).¹⁹ Initially, the FDA should finalize guidance on SaMD pre-specifications and algorithm change protocol inputs that FDA should periodically receive.

- Develop methods to efficiently communicate when FDA has authorized a product developed with or that utilizes AI/ML, along with information on how it was developed, is maintained and performs, and aligns with the latest best practices and regulatory requirements that ensure patient safety using easily understood (e.g., infographics) and standardized formats. For example, where approval is required for the deployment of new solutions in the market, the FDA should provide information describing the datasets used to train the AI/ML software and what efforts are being taken to align with ethical standards and to mitigate data biases. This work should build on the recently released database of AI-enabled devices legally marketed in the United States from the FDA's Digital Health Center of Excellence.²⁰
- Serve as a coordinator and convenor of other U.S. federal agencies to ensure a harmonized approach to health AI/ML transparency across government.
- Build on its leadership to date within the International Medical Device Regulatory Forum (IMDRF), promote its approach to SaMD AI/ML to improve approaches to transparency internationally.
- Host recurring public events, in partnership with health AI/ML developers, patients, and providers, that feature the FDA Digital Health Center of Excellence's latest approaches and thinking, as well as demonstrations of AI/ML in healthcare today.

The Centers for Medicare and Medicaid Services (CMS) should:

- Continue to develop its understanding of medical AI/ML definitions, present-day and future AI/ML solutions, how AI/ML is changing the practice of medicine, and the future of AI/ML medical coding.
- Develop Medicare support mechanisms for the use of AI/ML by providers based on clinical validation, alignment with clinical decision-making processes familiar to providers, and high-quality clinical evidence.
- Build on support provided in the Medicare system for the use of health AI,²¹ develop easy to understand resources for Medicare beneficiaries that capture how AI/ML is being used in the Medicare system and what it means to patients. CMS should leverage its Advisory Panel on Outreach and Education²² to develop this messaging.

19 As the FDA has noted, new reporting mechanisms for a scalable AI/ML medical device reporting structure “may require additional statutory authority to implement fully”. Proposed Regulatory Framework for Modifications to Artificial Intelligence/Machine Learning (AI/ML)-Based Software as a Medical Device (SaMD) - Discussion Paper and Request for Feedback (Apr. 10, 2021) at 15. Available at <https://www.fda.gov/files/medical%20devices/published/US-FDA-Artificial-Intelligence-and-Machine-Learning-Discussion-Paper.pdf>.

20 <https://www.fda.gov/medical-devices/software-medical-device-samd/artificial-intelligence-and-machine-learning-aiml-enabled-medical-devices>. This FDA list currently provides key information such as submission number, device and company name, and date of marketing authorization of the device (510(k) clearance, granting of De Novo, or PMA approval).

21 For example, CMS already provides payment for CPT code 92229 (point-of-care diabetic retinopathy automated analysis and provides a diagnostic report using AI).

22 <https://www.cms.gov/Regulations-and-Guidance/Guidance/FACA/APOE>.

The Federal Trade Commission (FTC) should:

- Support ways to mitigate biases or other unfair outcomes from healthcare AI,²³ and, where appropriate, enforce against violations of key laws such as Section 5 of the FTC Act, which prohibits unfair or deceptive practices, where appropriate.

Accrediting and Licensing Bodies, and Medical Specialty Societies and Boards should:

- Develop medical standard of care and ethical guidelines to address emerging issues with the use of SaMD AI/ML in healthcare needed to advance the quadruple aim.
- Develop and disseminate guidance and education on the responsible deployment of SaMD AI, both generally and for specialty-specific uses.

Academic and Medical Education Institutions should:

- Develop and include curriculum that will advance understanding of and ability to use healthcare AI/ML solutions, which should be assisted by inclusion of non-clinicians, such as data scientists and engineers, as instructors. Ongoing training and continuing education should also advance understanding of the safe and effective use of AI/ML in healthcare delivery, addressing both its capabilities and limitations.
- Develop curriculum to advance understanding of data science research to help inform ethical bodies such as Institutional Review Boards (IRBs) that are reviewing protocols of clinical trials of AI-enabled medical devices.

²³ <https://www.ftc.gov/news-events/blogs/business-blog/2021/04/aiming-truth-fairness-equity-your-companys-use-ai>

Conclusion

CHI is pleased to present its recommendations on AI/ML transparency for the consideration of the healthcare ecosystem, policymakers, and others. We are committed to continued engagement with the digital health community writ large to realize the both the responsible deployment of AI/ML across healthcare and its immensely positive societal benefit.

CHI Health AI Roles & Interdependency Framework

ConnectedHealthInitiative



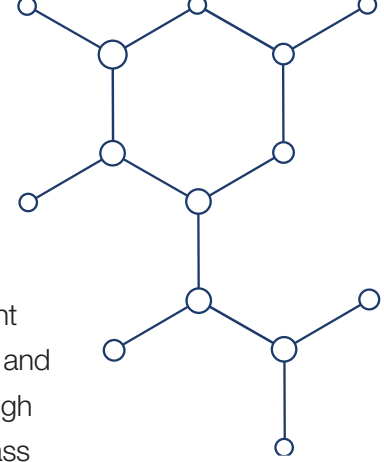
Overview

Artificial Intelligence (AI), especially generative AI, is already a powerful tool in healthcare, offering amazing potential to upgrade patient care by improving care outcomes and patient experiences, reducing healthcare provider burnout by simplifying administrative tasks, and helping to lower the total cost of care. One of the most helpful ways to see the value of AI in healthcare is to view the question through the lens of the “quadruple aim” framework. Built on the Institute for Healthcare Improvement’s “triple aim,” a widely accepted compass to optimize health system performance, the quadruple aim focuses on four key areas where health systems need to be improved, all of which AI is already, and will continue to, provide value across:

- Enhancing population health.
- Improving patient experience, satisfaction, and health outcomes.
- Augmenting clinician and healthcare team experience and satisfaction.
- Lowering overall costs of healthcare.

CHI has explored the ways in which AI is supporting each of the four aims of the quadruple aim in CHI’s paper, [Why Does Healthcare Need AI?](#)

But this promising technology is not infallible, and as healthcare organizations seek opportunities to use AI, stakeholders are facing important questions about how various risks or limitations should be handled in the development, distribution, deployment, and end use chain. Many organizations involved in the creation or application of healthcare AI have started to develop Responsible AI programs aimed at managing these risks or limitations within their organization. But as we have learned from other new technologies in the past, stakeholders can benefit from a clear discussion around all the safety measures and other actions that are needed, and how those actions might be applied at different steps from creation to the operation of the tool by the end user. This discussion will help various stakeholders better determine accountability for responsible AI best practices across this chain of stakeholders.



CHI urges all stakeholders in the healthcare ecosystem that are developing and using AI to align with **CHI's consensus health AI principles**, which recognize the shared responsibility for AI safety, efficacy, and transparency. CHI supports (1) leveraging a risk-based approach to AI harm mitigation where the level of review, assurance, and oversight is proportionate to potential harms and (2) those in the value chain with the ability to minimize risks based on their knowledge and ability, and having appropriate responsibilities and incentives to do so.

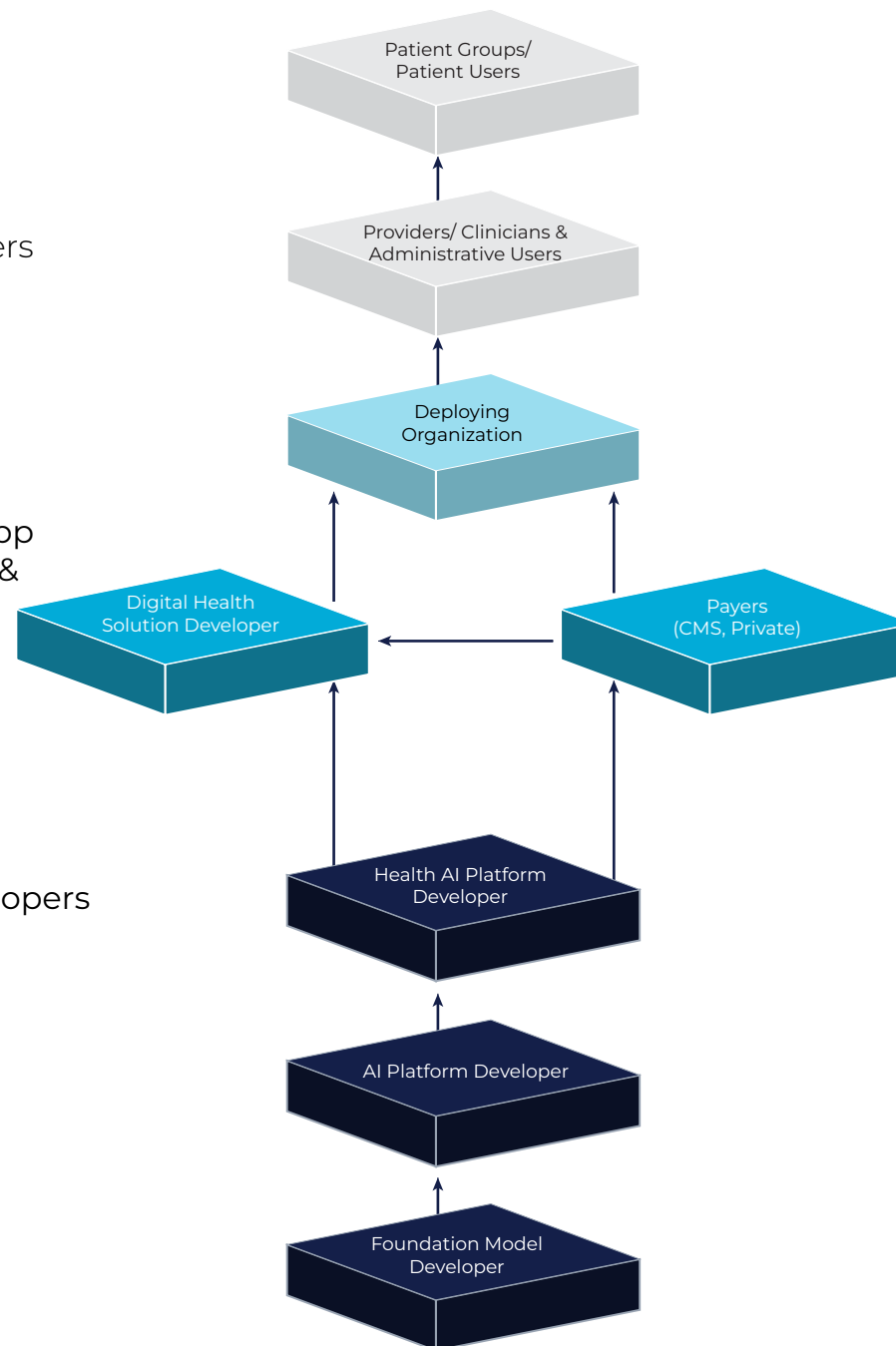
Further, managing AI/Machine Learning (ML) risks will be more challenging for small to medium-sized organizations, depending on their capabilities and resources. Building on these general health AI principles, CHI proposes clear definitions of stakeholders across the healthcare AI value chain, from development to distribution, deployment, and end use. Then, CHI suggests roles for supporting safety, ethical use, and fairness for each of these important stakeholder groups that are intended to illuminate the interdependencies between these actors, thus advancing the shared responsibility concept. These roles and interdependencies are also mapped to the Functions defined in the National Institute of Standards and Technology's (NIST's) AI Risk Management Framework (RMF).



1
Solution Users

2
Solutions/App
Developers &
Deployers

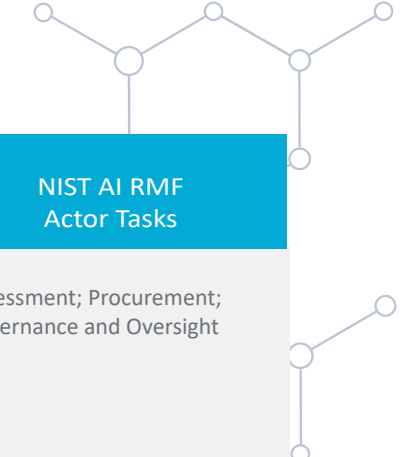
3
AI/ML Developers



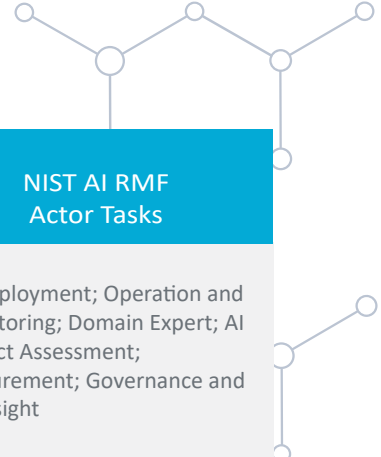
Note: Depending on the use case, some of the roles in the healthcare AI/ML value chain may be occupied by the same party; in other scenarios, some roles may not be occupied.

Stakeholder Group	Definition	Roles	NIST AI RMF Actor Tasks
AI/ML Developers	Someone who designs, codes, researches, or produces an AI/ML system or platform for internal use or for use by a third party. See below for defined Subgroups of this Stakeholder Group along with recommendations specific to that Subgroup.	- Informing deployers and users of data requirements/definitions, intended use cases/populations and applications (e.g., disclosing sufficient detail allowing providers to determine when an AI-enabled tool should reasonably apply to the individual they are treating), including whether the AI/ML tools are intended to augment human work versus automate workflows, and status of/compliance with all applicable legal and regulatory requirements. - Prioritizing safety, efficaciousness, transparency, and data privacy and security from the earliest stages of design, leveraging (and, where appropriate updating) existing medical AI/ML guidelines on research and ethics, leading standards, and other resources as appropriate. - Employing algorithms that produce repeatable results and, when feasible, are auditable, and make decisions that (when applied to medical care) are clinically validated, fostering efficacy through continuous monitoring. - Utilizing risk management approaches that scale to the potential likely harms posed in intended use scenarios to support safety, protect privacy and security, avoid harmful outcomes due to bias, etc. - Providing information that enables those further down the value chain can assess the quality, performance, and utility of AI/ML tools. - Aligning with relevant ethical obligations and international conventions on human rights and supporting the development of new ethical guidelines to address emerging issues as needed.	AI Deployment; Operation and Monitoring; Test, Evaluation, Verification, and Validation (TEVV); Human Factors; Domain Expert; AI Impact Assessment; Governance and Oversight

Stakeholder SubGroup	Definition	Roles
Foundation Model Developer	Someone who creates or modifies large and generalizable machine learning models that can be used/adapted for various downstream tasks and applications, such as natural language processing, computer vision, or software development.	Building on the cross-AI/ML Developer roles noted above: - Assessing what bias and safety issues might be present in its Foundation Model, and documenting steps taken to mitigate those issues in its Transparency Documentation (e.g., Transparency Notes, System Cards and product documentation). - Providing clear guidance on (1) how to use and adapt its Foundation Model for various foreseeable downstream tasks and applications, and (2) what limitations or risks may arise from doing so based on challenges discovered during testing and deployment.
AI Platform Developer	Someone who leverages existing foundation models and builds an industry-agnostic platform that enables other developers to access, customize, and deploy these models for various use cases and applications, such as natural language processing, computer vision, and/or software development.	Building on the cross-AI/ML Developer roles noted above: Testing for, identifying, and mitigating bias and safety issues that may arise from using or modifying existing foundation models for its AI Platform, and documenting these issues and steps taken to address them in its transparency documentation (e.g., transparency notes, system cards and product documentation).
Health AI Platform Developer	Someone who creates or uses AI-powered platforms that are tailored for the healthcare domain, such as administrative efficiency, diagnostics, therapeutics, or research. These platforms may leverage foundation models (or other types of machine learning models or solutions), such as AI platforms, that are suitable for specific healthcare problems and data sources.	Building on the cross-AI/ML Developer roles noted above: - Meeting specific requirements and standards of the healthcare domain, such as accuracy, efficacy, explainability, and compliance with regulations. - Testing for, identifying, and mitigating any bias and safety issues that may affect the health outcomes of patients or the performance of clinicians using the Health AI Platform, and documenting these issues and the steps it has taken to address them in its transparency documentation (e.g., transparency notes, system cards and product documentation).
Digital Health Solution Developer	Someone who creates complete digital tools and technologies to improve health and healthcare outcomes, such as providing diagnostic and administrative solutions for clinicians, patients, and healthcare organizations. They may build digital health solutions with both health AI platforms, which are specialized for the health care domain, and AI platforms, which are more general and adaptable for various use cases and applications.	Building on the cross-AI/ML Developer roles noted above: - Specifying appropriate uses for its digital health solution to avoid amplifying bias or safety issues that may exist in the underlying foundation models, AI platforms, or health AI platforms. - Designing user interfaces to enable an end user to safely and effectively act upon the output of the tool, such as providing explanations, feedback mechanisms, or human oversight options, providing clear documentation to Deploying Organizations and Users to help them avoid bias and safety issues.



Stakeholder Group	Definition	Roles	NIST AI RMF Actor Tasks
Deploying Organization (Healthcare Provider or Payor)	Someone who is a healthcare providers and health care payors that and is deploying solutions built by Digital Health Solution Developers. They may also have their own internal IT staff that use health AI platforms or general AI platforms to develop their own custom digital health solutions.	Respecting that managing AI/ML risks will be more challenging for small to medium-sized organizations depending on their capabilities and resources: - Adopting AI/ML Developer instructions for use, specifying appropriate uses for Users through governance policies to avoid bias and safety issues that may exist in the underlying foundation models, AI platforms, or health AI platforms. - Developing and leveraging digital health solutions that augment efficiencies in coverage and payment automation, facilitate administrative simplification/reduce workflow burdens, and are fit for purpose. - Setting organization policy/designing workflows to reduce the likelihood that a User will act upon the output of the tool in a way that would cause fairness/bias or safety issues (tailored explanations, feedback mechanisms, and/or human oversight options). - Developing and organizational guidance on how the digital health solution should and should not be used. - Creating risk-based, tailored communications and engagement plans to enable easily understood explains to patients about how the digital health solution was developed, its performance and maintenance, and how it aligns with the latest best practices and regulatory requirements.	Assessment; Procurement; Governance and Oversight
Provider/Clinician Users and Administrative Users	Someone who directly interacts with or benefits from the digital health solutions that are built by Digital Health Solution Developers or by the internal IT staff of the Deploying Organization. They may include clinicians, such as doctors, nurses, or pharmacists, and administrative staff, such as billing, claims, or customer service personnel, in the provider and payor organizations.	Respecting that managing AI/ML risks will be more challenging for small to medium-sized organizations depending on their capabilities and resources: - Taking required training and incorporating employer guidance about use of AI/ML digital health solutions. - Documenting (through automated processes or otherwise) whether AI is being used in medical records and report any issues or feedback to the developer, such as errors, vulnerabilities, biases, or harms (where AI/ML's use is known by the User). - Ensuring there is appropriate clinician review and review of the output or recommendations from each digital health solution prior to acting on it (where AI/ML's use is known by the User).	AI Deployment; Operation and Monitoring; Domain Expert; AI Impact Assessment; Procurement; Governance and Oversight



Stakeholder Group	Definition	Roles	NIST AI RMF Actor Tasks
Payer Users (Centers for Medicare and Medicaid Services [CMS], State Medicaid, Private)	Someone that pays for the cost of healthcare services administered by a healthcare provider.	- Leveraging AI/ML systems that improve efficiencies in coverage and payment automation, facilitate administrative simplification, and reduce provider workflow burdens. - Aligning with medical AI/ML definitions, present-day and future AI/ML solutions, the future of AI/ML medical coding changes and trends. - Developing support mechanisms for the use of AI/ML by providers based on clinical validation, aligning with clinical decision-making processes familiar to providers, and high-quality clinical evidence. - Assuring that AI/ML systems allow for the individualized assessment of specific medical and social circumstances and provider flexibility to override automated decisions, ensuring that use of AI/ML does not improperly reduce or withhold care, or overrides the provider's clinical judgement. - Disclosing information about training and reference data to demonstrate that AI/ML systems do not create or exacerbate inequities and that protections are in place to mitigate bias. - Developing and proliferating easy to understand resources for beneficiaries and their providers that capture how and when AI/ML is being used, what information it is leveraging, and what it means to patients.	AI Deployment; Operation and Monitoring; Domain Expert; AI Impact Assessment; Procurement; Governance and Oversight
Patient Groups/ Patient Users	Someone who uses digital tools and technologies that are built by Digital Health Solution Developers or experiences their use in treatment.	- Developing and proliferating easy to understand resources that capture how AI/ML is being used and what it means to patients/patient groups, including explanations on the purpose and limitations of the digital health solutions that they use or benefit from (e.g., diagnostic, therapeutic, administrative). - Raising awareness of patients' rights and choices when using digital health solutions, such as consent, access, correction, or deletion of their personal data.	Human Factors
Standard-Setting Organizations	An organization whose primary function is developing, coordinating, promulgating, revising, amending, reissuing, interpreting, or otherwise contributing to the usefulness of technical standards to those who employ them.	Developing and promoting adoption of international voluntary/non-regulatory consensus standardized approaches and resources to steward a shared responsibility approach to AI.	Human Factors; Domain Expert; AI Impact Assessment; Governance and Oversight

Stakeholder Group	Definition	Roles	NIST AI RMF Activities
Certification Bodies & Test Beds	A certification body is a third-party organization that assures the conformity of a product, process or service to specified requirements. A test bed is a platform for conducting rigorous, transparent, and replicable testing of scientific theories, computing tools, and new technologies to a standard.	- Creating and making available transparent and reliable processes for the assurance of conformity to voluntary AI standards. - Creating and making available voluntary sandbox environments to help evaluate the usability and performance of AI/ML-based high-performance computing applications to advance the understanding of how reliable and efficacious AI, and to provide an appropriate assurance of reliability and efficacy.	Test, Evaluation, Verification, and Validation (TEVV); Human Factors; Domain Expert; AI Impact Assessment; Governance and Oversight
Accrediting and Licensing Bodies, and Medical Specialty Societies and Boards	Accrediting and licensing bodies are governing authorities that establish the suitability of any participating certification body. Notably, state-level board serve this purpose for physicians, nurses, and other clinicians to standards set by each state. Medical specialty societies are organizations for physicians, research and clinical scientists who are actively involved in the study of a particular specialty.	- Based on clinical needs and expertise, developing and setting the medical standard of care and ethical guidelines to address emerging issues with the use of AI/ML in healthcare needed to advance the quadruple aim. - Identifying the most appropriate uses of AI-enabled technologies and developing and disseminating guidance and education on the responsible deployment of AI/ML in healthcare, both generally and for specialty-specific uses.	Test, Evaluation, Verification, and Validation (TEVV); Human Factors; Domain Expert; AI Impact Assessment; Governance and Oversight
Academic and Medical Education Institutions	Tertiary educational institutions, professional schools, or forms a part of such institutions, that teach medicine and awards a professional degree for physicians or other clinicians.	- Developing and teaching curriculum that will advance understanding of and ability to use healthcare AI/ML solutions responsibly, which should be assisted by inclusion of non-clinicians such as data scientists and engineers as instructors. - Developing curriculum to advance the understanding of data science research to help inform ethical bodies (e.g., Institutional Review Boards that are reviewing protocols of clinical trials of AI/ML-enabled medical devices).	Human Factors; Domain Expert; AI Impact Assessment