

ConnectedHealthInitiative

July 10, 2026

Matthew Diamond
Food and Drug Administration (FDA)
United States of America
International Medical Device Regulators
Forum
Artificial Intelligence/Machine Learning-
enabled
Working Group Co-Chair

Gabriela Commatteo
Medicines and Healthcare products
Regulatory Agency (MHRA)
United Kingdom
International Medical Device Regulators
Forum
Artificial Intelligence/Machine Learning-
enabled
Working Group Co-Chair

RE: *Connected Health Initiative Comments to the International Medical Device Regulators Forum on its Artificial Intelligence/Machine Learning-enabled Working Group's proposed "Technical Framework for Artificial Intelligence Life Cycle Management" (IMDRF/AIML WG/N93 DRAFT: 2026)*

The Connected Health Initiative (CHI) appreciates the opportunity to provide input to the International Medical Device Regulators Forum (IMDRF) as its Artificial Intelligence/Machine Learning-enabled Working Group develops the proposed Technical Framework for Artificial Intelligence Life Cycle Management (the "Framework"). We commend the Working Group's continued leadership in building globally harmonized approaches to AI-enabled medical devices, and we offer the comments below to help ensure the Framework remains risk-based, technology-neutral, and reflective of how AI-enabled medical devices are developed and deployed today.

CHI is a leading multistakeholder policy and legal advocacy effort driven by a consensus of thought leaders from across the connected health ecosystem. We aim to realize an environment in which patients and providers can improve health outcomes and reduce costs through connected health technologies. As part of its commitment to responsibly advancing AI in healthcare, CHI assembled a Health AI Task Force of innovators and experts that has developed a range of recommendations for policymakers, several of which we offer as contributions to IMDRF's work below. For more information on CHI, please visit <https://connectedhi.com/>.

Leveraging health data with AI holds tremendous promise for advancing value-based care across research, health administration and operations, population health, practice-delivery improvement, and direct clinical care. Realizing that promise requires a balanced approach that pairs innovation with safeguards calibrated to risk. We continue to urge IMDRF to evaluate AI in healthcare through the "quadruple aim" framework (enhancing population health; improving patient experience, satisfaction, and outcomes; improving the clinician and care-team experience; and lowering the overall cost of care) and to ground its expectations in the risk of harm and benefit associated with a given intended use.

CHI is broadly supportive of the Framework. In particular, we welcome its risk-based, context-dependent orientation, its recognition that responsibilities are shared across the AI value chain, and

its nuanced treatment of trade-offs—for example, the discussion in Section 5.3 of the trade-offs among model explainability, complexity, and performance, which appropriately ties expectations to the clinical context of use. To keep the document internally consistent with that risk-based approach, and to reflect current development practices such as the growing use of third-party foundation models and cloud architectures, we respectfully offer the following recommendations.

1. Preserve architecture and deployment neutrality in cybersecurity expectations (Section 4.4)

The cybersecurity discussion in Section 4.4 (lines 343–353) implies that cloud connectivity inherently introduces greater cyber risk and points to local hosting “if possible” as a mitigation. This framing is not accurate. A well-architected cloud deployment can deliver security equal to or greater than local hosting—for example, through centralized patching and monitoring, hardened key management, network isolation, and continuous threat detection. Whether a device is hosted locally, on-premises, or in the cloud, the appropriate question is whether the design achieves the required security outcomes.

We recommend that IMDRF adopt an architecture-neutral posture and revise these lines so that cybersecurity expectations are met through appropriate, risk-based design regardless of deployment model, rather than by suggesting a preference for any particular hosting arrangement.

2. Focus validation on the device, and calibrate training-data transparency to modern development practices (Sections 5.2, 5.3, and 5.4)

Section 5.2 treats detailed training-data transparency and characterization as a baseline expectation. In practice, manufacturers increasingly build AI-enabled medical devices on top of third-party foundation models and other general-purpose models for which training-data details are proprietary, incomplete, or simply unavailable. The Framework’s emphasis should therefore rest on the validation and testing of the finished device in its intended use, not on upstream training data. Where training-data information is limited, that limitation should inform the rigor and design of the manufacturer’s validation approach—it should not, by itself, preclude use of an otherwise suitable model. Relatedly, while CHI supports appropriate transparency about a model’s intended use, known limitations, and responsible-design choices, the expectation at lines 690–692 that “transparent information” be provided about general-purpose models can be difficult to satisfy for manufacturers incorporating third-party models, who may lack access to that information; transparency expectations should not require disclosure of proprietary third-party model or training data that the device manufacturer does not possess.

Section 5.4.1 frames “external validation of the model” as a distinct unit of analysis. We recommend clarifying that validation is properly conducted at the level of the AI-enabled medical device, assessed in its intended use environment and workflow, and that model-level external validation (line 746) is one supporting input rather than the endpoint. Anchoring validation at the device level is consistent with the Framework’s own recognition (in Section 5.4) that clinical evaluation covers the entire device rather than the model in isolation.

3. Keep monitoring, logging, and human oversight risk-proportionate and internally consistent (Sections 4.3, 4.4, 5.1, and 5.6)

The Framework is strongest where it is explicitly risk-based and context-dependent, as in the explainability trade-off discussion at lines 613–627. Several passages, however, call for “continuous” monitoring of all AI-enabled medical devices—for example at lines 346, 452–454, and 983–986—which is in tension with that risk-based approach. By way of comparison, the U.S. Food and Drug Administration’s draft guidance on AI-enabled device software does not require performance-monitoring plans for every device (including many cleared through the 510(k) pathway). We recommend harmonizing the document so that the intensity and frequency of monitoring scale with device risk and context of use, and replacing blanket “continuous” expectations with language tied to risk.

For the same reason, the logging recommendations in Section 5.6 (lines 970–975)—which contemplate capturing model predictions, explainability elements, heatmaps, top predictive features, data inputs, model outputs, timestamps, user interactions, model versions, and environmental context—should be applied in a risk-proportionate manner rather than as a uniform expectation for all AI-enabled medical devices. The appropriate scope and granularity of logging should reflect the device’s risk profile and context of use.

Finally, Section 4.3 describes human oversight as “essential throughout the entire life cycle.” As written, this constrains autonomous and semi-autonomous devices, for which continuous human oversight may be neither necessary nor appropriate. We recommend that oversight expectations be framed as context- and risk-dependent, applied at the points in the life cycle where oversight meaningfully contributes to safety and effectiveness given the device’s context of use, rather than universally.

4. Clarify the allocation of responsibility across the AI value chain (Section 4.4)

The Framework at times blurs the responsibilities of AI developers and the manufacturers who deploy models in medical devices. The reference to “primary responsibility” at line 357, for instance, can be read to imply that third parties bear secondary regulatory obligations. To avoid confusion, we recommend that the Framework state clearly that responsibility for regulatory compliance rests with the legal manufacturer of the AI-enabled medical device, and that oversight and audit expectations placed on suppliers be bounded by what each entity can reasonably assess within its own scope of control. This approach is consistent with CHI’s Health AI Roles & Interdependency Framework, which articulates a shared-responsibility model across the value chain while keeping regulatory accountability with the manufacturer.

5. Strengthen and right-size guidance for generative AI (Section 5.6)

The Framework is largely oriented toward traditional machine learning, and its generative AI provisions (for example, lines 987–995 and 1019–1029) are comparatively underdeveloped and do not fully account for the differing risk profiles of these technologies. As one illustration, collecting interaction logs and prompt-response histories for all large language model-based devices would be burdensome in lower-risk contexts. Consistent with our recommendations above, we encourage IMDRF to expand its generative AI-specific guidance and to make the associated monitoring and

logging expectations risk-based—qualifying prompt- and interaction-log collection so that it applies where warranted by the risk and context of use.

6. Incorporate ISO/IEC 42001 as a foundational standard (Sections 3 and 4.1)

The Framework references ISO 13485 and a number of AI-related standards but does not reference ISO/IEC 42001:2023, the international standard for AI management systems. ISO/IEC 42001 is being adopted with increasing frequency across jurisdictions, including by leading AI service providers, and provides a management-system foundation well suited to the life-cycle governance the Framework describes. We recommend adding ISO/IEC 42001:2023 to the references in Section 3 and identifying it in Section 4.1 (Quality Management System) as a foundational AI management-systems standard for manufacturers to consider.

Supporting CHI resources

To assist the Working Group in its continued development of the Framework, we offer the following CHI resources:

- **CHI’s Health AI Policy Principles**, a set of recommendations on the range of issues policymakers should address in examining AI’s use in healthcare (<https://actonline.org/wp-content/uploads/Policy-Principles-for-AI.pdf>);
- **CHI’s Good Machine Learning Practices for FDA-Regulated AI**, a proposed risk-based approach addressing both locked and continuously-learning AI systems that meet the definition of a medical device (<https://connectedhi.com/wp-content/uploads/2022/04/CHIAITaskForceGMLPs.pdf>);
- **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
- **CHI’s Health AI Roles & Interdependency Framework**, which defines stakeholders across the healthcare AI value chain—from development to distribution, deployment, and end use—and proposes roles for supporting safety, ethical use, and fairness, illuminating the interdependencies among these actors and advancing the shared-responsibility concept (<https://connectedhi.com/wp-content/uploads/2024/02/CHI-Health-AI-Roles.pdf>).

CHI appreciates the opportunity to submit these comments and urges IMDRF's thoughtful consideration of the input above. We look forward to continuing to work with IMDRF and other stakeholders to realize the benefits of AI in healthcare.

Sincerely,

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

Brian Scarpelli
Executive Director

Chapin Gregor
Policy Counsel

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