

AyurDatta Solutions

Independent, physician-led safety governance for AI-assisted and real-time clinical trials

info@ayurdatta.com · www.ayurdatta.com

June 25, 2026

Via Electronic Submission — <https://www.regulations.gov>

Division of Dockets Management (HFA-305)
U.S. Food and Drug Administration
5630 Fishers Lane, Rm. 1061, Rockville, MD 20852

Re: Docket No. FDA-2026-N-4390 — “AI-Enabled Optimization of Early-Phase Clinical Trials Pilot Program; Request for Information” (91 FR 23100)

To the Division of Dockets Management and the Office of the Commissioner:

AyurDatta Inc. is an independent, physician-led clinical safety and governance firm focused on AI-assisted and real-time clinical trials in oncology and cell and gene therapy. We appreciate the opportunity to respond to the Agency’s request for information regarding the AI-Enabled Optimization of Early-Phase Clinical Trials Pilot Program. Our perspective is grounded in direct experience: our President and Chief Medical Officer is a board-certified oncologist who has provided safety leadership across 135 oncology trials, including 11 first-in-human studies, and has supported 27 INDs, 9 NDAs, and 3 BLAs, including a CD19 CAR-T cell therapy; our Chief Executive and Technology Officer architected a unified eClinical platform that has supported more than 1,000 trials.

We strongly support applying AI to the well-documented inefficiencies of early-phase development, and we agree that safety monitoring is among its highest-value applications. We write to offer one focused recommendation that cuts across the Agency’s questions on pilot design, evaluation metrics, and trustworthy AI: **every AI-assisted safety output in a clinical trial creates a matching obligation for documented, qualified human judgment.** A model may detect, score, and draft; it cannot adjudicate benefit-risk or bear regulatory accountability. We therefore recommend that the pilot explicitly require, structure, and measure the human-in-the-loop decision and the audit trail it produces — most urgently in the high-acuity oncology settings, such as the CAR-T and cell-therapy programs that already comprise the Agency’s proof-of-concept trials, where toxicities such as cytokine release syndrome and immune effector cell-associated neurotoxicity syndrome can escalate within hours.

Per the Agency’s instruction, our responses below are identified by their category and question number.

A. Pilot Program Design and Implementation

Question A.1 (Scope and Focus) — responding to 1.a and 1.c

We recommend prioritizing first-in-human and oncology dose-escalation trials, particularly in cell and gene therapy. These settings combine the greatest uncertainty, the smallest patient

populations, and the most acute, time-sensitive toxicities — the conditions under which both AI assistance and disciplined human oversight deliver the most value and carry the highest stakes. Among AI use cases, we recommend prioritizing safety monitoring, because it is where AI most directly intersects a non-delegable regulatory obligation: the benefit-risk decision.

Foregrounding safety monitoring lets the pilot test, in the highest-consequence setting, how AI assistance and human accountability operate together.

Question A.3 (Collaboration Models) — responding to 3.c

AI governance in a trial should be anchored to a named, credentialed physician accountable for adjudicating AI-flagged safety signals — not to a committee in the abstract, and not to the tool's developer. We recommend the pilot require each participant to designate, in advance, the qualified individual responsible for benefit-risk decisions on AI-assisted outputs, and to document that individual's credentials and decision authority. Investigators must retain clinical decision authority; the AI's role is to assemble and to suggest, never to decide. Clear human ownership is the precondition for meaningful governance.

Question A.4 (Operational Structure) — responding to 4.c

Participants will arrive with very different levels of safety-governance maturity, and the pilot should assess that dimension explicitly rather than assume it. We suggest a simple maturity scale for safety oversight, distinct from AI sophistication:

- (1) Absent — no documented process or owner;
- (2) Ad-hoc — work occurs but is person-dependent and reconstructable only after the fact;
- (3) Defined — a documented, owned process whose records nonetheless lag or scatter; and
- (4) Inspection-ready — self-evidencing records that are attributable, timestamped, and retrievable on demand.

A participant can be highly advanced in AI and still sit at level 2 in oversight maturity. Measuring this dimension will tell the Agency whether AI gains are being realized on a defensible foundation.

B. Evaluation Metrics and Success Criteria

Question B.2 (Decision Quality) — responding to 2.a and 2.b

Beyond concordance between AI-supported and traditional decisions, we recommend the pilot measure whether each AI-supported decision is accompanied by a documented, attributable human adjudication. Two practical metrics: a **documented-adjudication rate** (the share of AI-assisted safety decisions carrying a complete, signed physician review) and **time-to-documented-decision** (the interval from signal detection to a recorded benefit-risk disposition). Together these capture not only how fast and accurate the AI is, but whether the accountable human decision actually occurred and was recorded.

Question B.3 (Participant Safety and Data Integrity) — responding to 3.a and 3.b

For safety-signal detection and response, we recommend measuring **time to documented physician disposition** rather than time to flag alone; a faster flag without a faster documented decision does not improve patient safety. We further recommend tracking the **physician override rate** — how often the reviewer changes the AI's suggested grade, causality, or

disposition. A non-trivial override rate is direct evidence that human oversight is substantive rather than a rubber stamp; a rate trending toward zero may itself be a finding worth examining.

Question B.5 (Trustworthiness, Aligned With NIST AI RMF) — responding to 5.b and 5.c

5.b (safety and risk mitigation). We recommend the pilot evaluate whether a participant's architecture makes human oversight structurally unavoidable rather than merely encouraged. The strongest safeguard is an invariant: no AI output becomes a record of decision without a qualified physician's review and electronic signature, with the bypass engineered out of the workflow. Systems that permit automatic disposition of safety-relevant outputs should be treated as higher risk regardless of model quality.

5.c (transparency and explainability across sponsor-developed and proprietary systems).

The Agency rightly asks how to evaluate trustworthiness for proprietary systems whose internals cannot be fully inspected. We suggest relocating the evidence from the model's interior to the **decision boundary**. Regardless of whether the underlying model is open or proprietary, trustworthiness can be evidenced by a tamper-evident audit trail that captures, for every AI-assisted safety output: the model and version used; the inputs and outputs; the reviewing physician's identity and credentials; the criteria applied; any modification or override; and the final signed disposition with timestamp. This decision-boundary record is uniformly measurable across vendors and modalities, it is exactly what an inspection requires, and it does not depend on disclosing proprietary model details. We believe it offers the Agency a practical, vendor-neutral basis for a trustworthiness metric.

Question B.7 (Qualitative Outcomes) — responding to 7.a

In our experience, stakeholder trust in AI-enabled trials tracks closely with the visibility and defensibility of the human-accountability layer. A concrete proxy: can an investigator, sponsor, or inspector retrieve, on demand, the documented physician review for a given set of AI-flagged signals? When the answer is yes and immediate, trust follows; when it requires reconstruction, trust erodes. We recommend assessing the retrievability of the documented human decision as a measurable indicator of trust.

We commend the Agency for advancing this initiative and for foregrounding trustworthy AI and alignment with the NIST AI Risk Management Framework. AyurDatta would welcome the opportunity to serve as a resource to the Agency and to pilot participants on the medical-oversight and documentation dimensions discussed above. Thank you for your consideration.

Respectfully submitted,

Ashok Srivastava, MD, PhD, MBA

President & Chief Medical Officer, AyurDatta Solutions

Elias Tharakan

Chief Executive & Technology Officer, AyurDatta Solutions

Contact: info@ayurdatta.com