

A Call to the CHAI Community

Building the Trust Infrastructure for the Next Era of AI

Dear CHAI Community,

We are entering a consequential moment in the development of artificial intelligence.

Bill Gates recently wrote that AI could become either an extraordinary equalizer or a profound source of injustice, and argued that the choices we make now will shape which future we inherit. Across the world, a broader debate is emerging about whether the development of increasingly capable AI is moving faster than our ability to understand, govern, and safely deploy it.

That debate is crossing traditional boundaries. Policymakers in the United States are debating proposals to slow or restrict the development of advanced AI. Internationally, leaders are convening frontier AI companies to discuss how this technology can be developed in ways that benefit society.

Most strikingly, the conversation is happening inside the AI laboratories themselves.

Anthropic CEO Dario Amodei has called on the industry to “pace the frontier,” arguing that capability development should give safety, alignment, and evaluation practices time to catch up. His proposal includes embedding independent third-party evaluators inside frontier AI companies. Microsoft AI recently published a Humanist AI Code of Conduct grounded in a similarly straightforward premise: “people matter more than AI,” with an explicit commitment to keeping advanced systems subordinate to meaningful human control.

These developments should command our attention.

But for the CHAI community, I believe the appropriate response is neither panic nor passivity.

It is action.

AI is advancing at extraordinary speed across health care and the life sciences. We are moving from systems that generate information to systems that can reason, recommend, coordinate, and increasingly act.

Greater capability must be accompanied by greater confidence that these systems are safe, effective, secure, and serving the people they are intended to help.

That is especially true in medicine.

Healthcare is built on trust. Patients entrust physicians with intimate information, consequential decisions, and sometimes their lives. They expect their clinicians to be advocates who understand their values and wishes and place those interests at the center of care.

If AI is to truly transform health, it must become a trusted tool within that relationship, not something that undermines it.

And that is precisely why CHAI was created.

Since our inception as a nonprofit founded and led by clinicians, CHAI has functioned as a community committed to sharing real world experience and through that process technical specifications, best practices, and testing and evaluation frameworks needed to assess AI before deployment and monitor it afterward.

Our premise has always been that responsible AI cannot be established through principles alone. Principles must become practice. We need evidence about whether an AI system works for its intended use, whether it performs safely across populations and settings, what its limitations are, and what happens to its performance over time. CHAI's work now spans governance, validation, use-case-specific testing and evaluation, and post-deployment monitoring.

Just as importantly, CHAI has built a community capable of doing this work together.

Today, that community includes thousands of clinicians, health system business and technical leaders as well as nonprofits, researchers, public-sector leaders, startups, large technology firms and the major frontier AI laboratories.

That breadth is not incidental. It is CHAI's greatest strength.

No single developer should determine what constitutes sufficient evidence for its own technology. No individual health system should have to recreate the same evaluation from scratch. And no policymaker should be asked to govern rapidly changing technology without access to the practical experience of the clinicians, developers, researchers, and patients using it.

CHAI provides the neutral table around which those communities can share experience and expertise, and ultimately develop consensus.

And we can do so in ways that protect the things our members appropriately need to protect: beginning with patient private health data but also proprietary intellectual property and commercially sensitive information.

The goal is to agree on what evidence should be generated, how it should be evaluated, and what information decision-makers need to determine whether an AI system is worthy of trust.

This is also why I believe the current discussion about independent third-party evaluators represents an important moment for our CHAI community.

Frontier AI companies need evaluators who understand not only model-level capabilities, but the realities of specific industries. Healthcare is highly contextual. A model that appears safe in a laboratory may behave very differently when incorporated into a clinical workflow, connected to sensitive information, or given increasing autonomy.

CHAI can help create a health-specific methodology for independent evaluation—one grounded in industry-developed consensus best practices, use-case-specific testing and evaluation frameworks, strong privacy and security protections, and real-world evidence from health systems. This is something we have been doing since our coalition came together at the beginning of 2021.

Our work in actively developing a regulatory sandbox playbook with policymakers, particularly at the state-level is beginning to create environments where these evaluations can transparently happen.

This is not about slowing innovation for its own sake, it is about creating the conditions in which innovation can scale responsibly.

Health systems must help generate real-world evidence and monitor performance. Technology companies must participate in transparent, credible evaluation. Frontier labs must enable meaningful independent assessment. Policymakers must create environments that permit responsible experimentation. Researchers and nonprofits must help us ask the hard questions. And patients and communities must have a meaningful voice in defining what trustworthy AI looks like.

Over the coming weeks, the CHAI leadership team will convene conversations across these communities to understand where current safety and evaluation practices are strong, where gaps remain, and where shared standards could add value.

We will structure these discussions to protect confidentiality and proprietary information, beginning where appropriate with smaller or one-on-one conversations before bringing the broader community together around areas of emerging consensus.

If your organization would like to participate in this work, I want to hear from you!

This is not a moment for any one organization to claim that it has all the answers.

It is a moment for a community to build them together.

CHAI has spent years creating the relationships, technical frameworks, consensus processes, and trust needed to do exactly that. Our responsibility now is to put that infrastructure to work at the scale and speed this moment demands.

I remain profoundly optimistic about what AI can mean for health.

Our opportunity is not simply to build more powerful AI. It is to ensure that the most powerful AI ever created is safe, worthy of public trust, respectful of human dignity, and ultimately directed toward human flourishing.

That is work worth doing—and it is work we can only do together.

A handwritten signature in dark ink that reads "Brian Anderson, M.D." The signature is written in a cursive, flowing style.

Brian Anderson, MD

Co-founder and Chief Executive Officer
Coalition for Health AI

Endnotes

1. Bill Gates, "A turbulent AI era and critical choices to make," Gates Notes, 2026.
2. Bernie Sanders and Greg Casar, "Sanders, Casar Introduce Legislation to Ban Artificial Superintelligence and Temporarily Pause Advanced AI Development," September 2026; and reporting on international discussions among frontier AI leaders in Scotland, September 2026.
3. Dario Amodei, "We Must Pace the Frontier," September 2026. Amodei proposes pacing frontier capability development and greater use of independent third-party evaluation.
4. Microsoft AI, "Humanist AI Code of Conduct," 2026.
5. Coalition for Health AI, 2026 program and governance materials. CHAI work includes AI governance playbooks, use-case workgroups, validation, testing and evaluation, and post-deployment monitoring.
6. Coalition for Health AI, 2026 community and membership materials. Community figures cited in this letter reflect CHAI's current organizational counts.

7. Coalition for Health AI, Regulatory Sandbox Policy Work, 2026. CHAI is developing a sandbox playbook and state implementation guide, partnering with NCSL, participating in Utah sandbox work, and engaging domestic and international partners.

8. Coalition for Health AI, Public Health AI Leadership Council (PULSE), 2026. PULSE supports public-health agencies with enterprise AI environments and real-world use-case pilots.

9. Coalition for Health AI, “Ensuring Health AI Supports Human Flourishing,” January 2025; CHAI 2026 workgroup materials describing work to align generative-AI use cases with human values.