

BOARD ROOM ESSAY

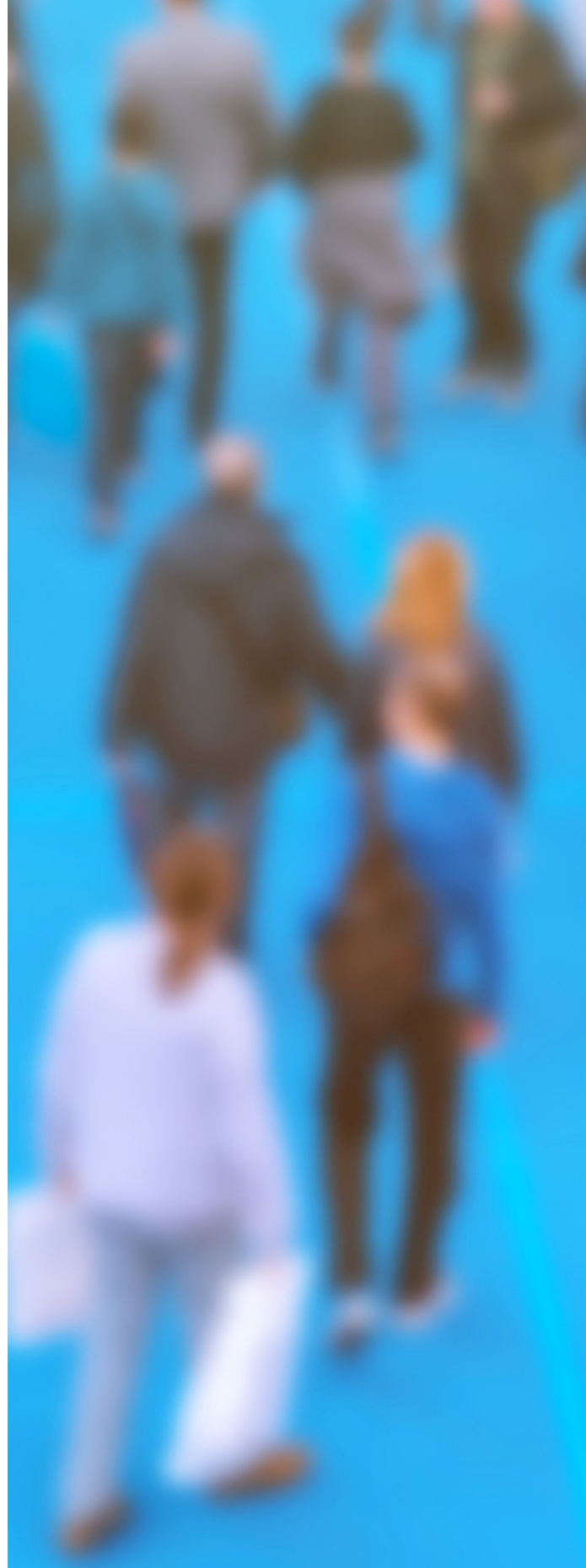
EXECUTIVE BRIEF NO. 004

Three Questions About Responsible AI

An AI system can be legal, compliant, and technically accurate — and still be wrong. Three questions reveal whether your organization is ready to be accountable for the answer.

AI GOVERNANCE

5 MIN READ • AUGUST 11, 2026



Three Questions Every Leader Should Ask About Responsible AI

The greatest risk of AI in health and medicine isn't that it fails. It's that it works — and in working, it scales. Three questions reveal whether you are ready to be accountable.

BY TOM LAWRY

The greatest risk of AI in health and medicine isn't that it fails.

It's that it works.

And in working, it scales. And in scaling, it quietly encodes the inequities of the real world into the intelligent one.

A system that improves outcomes overall while systematically underperforming for specific populations is not a breakthrough. It is an equity blind spot at scale.

Responsible AI demands that leaders confront this reality before deployment, not after. The question is not whether your AI works. It is who it works for, how it works, and whether you are willing to be accountable for the answer.

Done right, AI is not just about optimizing outcomes. It's about aligning outcomes with your personal and organizational values.

What Responsible AI actually means

Strip away the jargon and Responsible AI is simple to state: it is the practice of designing, deploying, and governing AI systems so they are safe, fair, transparent, and accountable — so the benefits of intelligence land equitably across every population you serve, not just on average.

It is not a technical standard to be met once and filed away.

It is an ongoing organizational commitment. And like every commitment that matters, it belongs to leadership — not to the IT department, and not to the vendors.

"An AI system can be legal and compliant. And it can still be wrong."

The gray zone between legal and right

AI is increasingly shaping clinical decisions — sometimes directly, often indirectly. The systems making those decisions are trained on data that reflects the world as it is, not as we wish it to be. That world includes decades of disparities in access, treatment, and outcomes.

Left unexamined, AI doesn't eliminate those disparities. It learns them. It encodes them. And, if left unchecked, it scales and amplifies them.

Meanwhile, the rules are still being written. The European Union has its AI Act. The American Medical Association, and others have proposed voluntary frameworks for the United States. But there is — and will continue to be — a gap between what technology can do and how it is governed.

That gap creates a gray zone. And in that gray zone, something important is true:

An AI system can be legal and compliant. And it can still be wrong.

That is why Responsible AI cannot be outsourced to regulators. The right approach is not to wait for perfect standards. It is to act — with intention — guided by principles that reflect your organization’s mission and values.

The gray zone between what is legal and what is right will not disappear. It will only grow. That is not a reason for hesitation.

It is a call for leadership.

Three questions that reveal where you stand

A Responsible AI framework is only as valuable as its application. These three questions give leadership teams a practical starting point to assess where your organization stands — and where attention is needed now.

Question 1: Have you formally adopted a Responsible AI framework?

Formally is the operative word. A framework becomes real when it is reviewed and approved at the highest levels of the organization — management and the Board. A draft circulating among IT staff is not an adopted framework. The test is whether your organization has made an explicit, documented commitment.

A sound framework addresses six core principles:

- **Fairness** — AI systems treat all patients and populations equitably
- **Reliability & Safety** — systems perform consistently and without causing harm
- **Privacy & Security** — sensitive health data is protected
- **Inclusiveness** — AI works well for every population served
- **Transparency** — systems and their outputs can be understood

- **Accountability** — people, not algorithms, remain responsible for AI-driven decisions

Formal adoption is how an organization commits to asking the harder question: not just whether AI works, but who it works for.

Question 2: Do you apply it to every AI application — before and after deployment?

Adoption is only the first step. A framework that is rarely consulted — or applied unevenly across departments — offers little protection, and may even add legal and reputational risk rather than reducing it.

Consistent application means evaluating every AI tool at two moments. Before deployment, when design choices can still be changed. And continuously after, because model performance drifts — in ways that are invisible without active monitoring.

Done consistently, this isn’t bureaucracy. It’s how AI tools stay both effective and equitable across every population, not just on average.

Question 3: Does vendor procurement require alignment with your framework?

Most of the clinical AI your organization deploys won’t be built in-house. It will arrive embedded in the products you purchase and license — an EHR upgrade, a radiology module, a clinical decision support tool. Each carries its own model assumptions, its own training data, and its own potential biases.

Your framework has to travel with your money. Build its requirements explicitly into your RFI and RFP processes, and make alignment a condition of purchase.

Because here is the reality no contract changes: when you deploy a vendor’s AI, your patients experience its outcomes. And your organization owns the accountability.

THREE QUESTIONS FOR YOUR NEXT LEADERSHIP AGENDA

- Have you formally adopted a Responsible AI framework?
- Is it applied to every AI application — before and after deployment?
- Does vendor procurement require alignment with your framework?
- The goal is not a perfect score. It is to surface the gaps — and close them.

The executive suite and board room conversations

The goal is not a perfect score on three questions. It is to surface the gaps — and begin working to close the gaps you have defined.

Prediction: Within five years, Responsible AI governance will be scrutinized the way financial controls are today — by boards, by regulators, and by the public. The organizations that can show their work will be trusted with more.

Recommendation: Put these three questions on your next leadership agenda. Answer them honestly. Where the answer is no, assign an owner and a date. If you need scaffolding, start with the NIST AI Risk Management Framework and the Coalition for Health AI's blueprint for trustworthy AI at chai.org — then adapt them to your mission and values.

Innovation does not scale without trust. Responsible AI is how you build both.

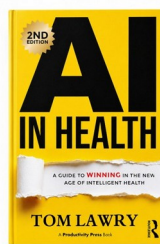
T.

SOURCE

NIST AI Risk Management Framework, National Institute of Standards and Technology (2023); Blueprint for Trustworthy AI, Coalition for Health AI (2023), www.chai.org.

Tom Lawry is managing director of Second Century Tech, author of [AI in Health](#), and the former national director of AI for health and life sciences at Microsoft.

www.tomlawry.com



Available in October from all major booksellers.

[Preorder from Amazon](#)

© 2026 Tom Lawry. All rights reserved.

This essay reflects the author's views and is for general information only; it is not legal or clinical advice.