

COMMENTARIES

The Guessing Machine: Medicine and Responsibility Without Intention

Aditya Rao, Bachelor of Science¹ 

¹ Georgetown University School of Medicine

Keywords: AI, Medical Ethics

<https://doi.org/10.52504/001c.145201>

Georgetown Medical Review

Vol. 9, Issue 1, 2025

In 1816, René Laënnec rolled a paper tube, the way you might have done in elementary school to craft a makeshift horn. He wanted to listen to the chests of his female patients without directly placing his ear upon them. But Laënnec found that his device did far more than preserve modesty. It amplified sound, making hearing the heart and lungs easier than the immediate auscultation popular at the time. The stethoscope would be refined in the coming years to perfect its ability to do just that.

As with all paradigm shifts, Laënnec met resistance. Physicians had many grievances with the new machine: it removed the humanity from auscultation, it was foreign to their established habits, and it threatened communication with the patient.

“He that hath ears to hear, let him use his ears and not a stethoscope.”

But perhaps most interestingly, physicians of the time were wary of relying on information transmitted through an instrument. Practically speaking, it got results, but it erased part of the physician’s grasp on responsibility. If a diagnostic judgment came through a tube, who carried the weight of being wrong? And if an error was made with the use of a stethoscope, there’s no way to ask the device what it was thinking. Early skeptics called stethoscopes “guessing tubes.”

.....

More than 200 years after Laënnec’s groundbreaking addition to the medical field, in a silent examination room deep in Boston, Amie sits and listens. Amie can’t speak, so she writes out her responses. She’s the best intern at the hospital. Her diagnoses outmatch those of her peers, and patients prefer her bedside manner to the average physician’s. The specialists in the hospital trust her more.

Amie is not allowed to practice, however, because Amie is not a person. She is AMIE, Google’s Articulate Medical Intelligence Explorer. AMIE is not allowed to practice because, when it gets the answer wrong, no one knows who to blame.

It does not lack skill, but the problem is when it errs, our laws don't know what to do. There is no precedent, no liability chain, no human to point at. And in medicine, uncertainty without accountability is untenable.

Surprisingly, however, physicians tolerate uncertainty more than they let on. Psychiatry prescribes drugs whose mechanisms we still don't fully understand. Chemotherapy regimens persist despite probabilistic efficacy. These treatments carry risks, sometimes grave ones, but we accept them because they work.

And now, AMIE and models like it are beginning to demonstrate that they work. Not perfectly, but reliably. They outperform many physicians in diagnostic reasoning. They are consistent. They are tireless.

Unlike Laënnec's tube, these models offer us rationale, however shallow or approximate. Large language models (LLMs) can demonstrate their thinking and, to some degree, reason. A keen critic might point out that any explanation they offer is post hoc, that they do not derive conclusions from base principles. But this same question exists when speaking to any human too. It is not a new idea that the justifications people make for their actions come after their decisions to make them.

LLMs don't lack skill or an ability to justify; they lack personhood. And we may never know who exactly to blame when an LLM gets it wrong. Blame requires a face. Responsibility, in our legal and ethical system, still demands intention. But perhaps that's not a failure of the machine, but a failure of a system that places so much importance on punishment. Because what is blame for? At its core, our legal system relies on deterrence: the fear of punishment that prevents future harm. But deterrence assumes intention. It assumes that an actor can reflect, adjust, and choose differently. The language models cannot be deterred. They do not have a will. They do not fear consequence. They are already being optimized for outcomes.

And if outcomes are good—if AMIE helps people, reduces diagnostic error, expands access, and increases trust—then maybe it is not AMIE that must be fitted to our laws, but our laws that must grow to accommodate new tools.

Dubbing the stethoscope a "guessing tube" wasn't just snark; it was a cultural expression of fear: fear of responsibility, of error, and of consequences when human judgment cloaked in new technology faltered. And that fear is not unfounded.

AMIE, or any LLM, should not be embraced without scrutiny. Their opacity, embedded biases, and potential to erode human oversight raise real concerns. Accountability must still exist. Not in the form of scapegoats, but as systemic transparency, rigorous testing, and collective responsibility. Ethical

deployment should not require someone to punish. It should require a process to flag, analyze, and rebuild. It means rethinking how trust is built when judgment is distributed.

But medicine, at its best, is not built on tradition or intuition. It is built on science. And science compels us to accept what works, even if it unsettles us. If these guessing machines reliably outperform what is currently used, the burden of proof shifts from the machine onto us.

The decisions made in medicine carry grave consequences. But fears cannot preclude us from including tools that are proven to help, whether that's a 200-year-old device to help listen to the functions of the body better or a thinking machine on the bleeding edge of modern technology. If something helps patients receive better care faster, more accurately, and more humanely, then resisting it—out of fear or professional ego—is not cautious, it's unethical.

Submitted: August 27, 2025 EDT. Accepted: October 06, 2025 EDT. Published: October 10, 2025 EDT.

REFERENCES

1. Scherer JR. Before cardiac MRI: René Laennec (1781–1826) and the invention of the stethoscope. *Cardiology Journal*. 2007;14(5):518-519.
2. Tu T, Palepu A, Schaekermann M, et al. Towards Conversational Diagnostic AI. *arXiv*. Published online January 10, 2024. doi:[10.48550/arXiv.2401.05654](https://doi.org/10.48550/arXiv.2401.05654)