

An Executive Field Guide to Implementing Trustworthy AI.

Information → Influence → Intention (Triple III Model)

AI must earn trust before it earns autonomy

By Dr. James L. Norrie, DPM, LL.M | October 13, 2025

Series Overview

We previously outlined our **SAFER AI** framework for ethical and responsible systems design. This series moves from principle to practice: how to make AI trustworthy in the real world and how to convert good intent into worthy action at the moment of human decision. A follow-on series will detail strategic use cases drawn from pilot results.

Experience suggests most enterprise teams try to “fix” AI in the wrong ways, something we see all the time. Teams chase better content—more data, faster models, slicker workflows—then wonder why people still do not rely on the system when it counts most. That is psychologically ignorant of basic scientific facts about how humans gather information and make decisions. Our view is therefore very different: higher human trust in AI is the missing ingredient. To close that gap, we focus on context—who the person is, how they decide, and what it takes for an agent to earn trust before exercising any autonomy, especially when stakes are high. When content and context intersect at the point of query and reply, we proved AI output becomes human outcomes.

Why This Works (and Why Others Stall)

In our view, human delegation to agentic AI cannot be assumed. Or forced as a default setting because trust rapidly erodes as users detect unexpected gaps. Instead, autonomy in our model is earned through five specific, sequential trust gates—Provenance, Fit (**myQ®**), Stakes, Reversibility, and Ethical Alignment—which determine if and how an agent should act. Clear alignment with high confidence permits limited autonomy. Ambiguous alignment downgrades autonomy and requires explicit confirmation. Failed alignment declines the move and explains why. At every rung the human can slow, pause, or roll back. If any gate fails, the system reverts to counsel, not control. Autonomy becomes a privilege earned by performance under constraint.

Most solutions optimize what a model perfunctorily says. Triple III optimizes when to say it, how to say it, and whether to act at all. That is why it improves compliance, reduces accidental-insider risk, and raises adherence in finance, healthcare, education, and other domains where trust and influence are inseparable and outcomes are measured in avoided errors and on-time follow-through. Early pilots show that when trust and fit rise together, people are more willing to consider, comply with, and continue safer, higher-quality actions.

Understanding the Triple III Model

Our patented Triple III Model operationalizes this shift through a staged, testable flow: Information → Influence → Intention. First, the system earns trust with verifiable evidence. Only then does it adapt how it communicates to fit the person in front of it. Finally, it converts guidance into a safeguarded plan with preview, confirm, and undo. Gates sit between steps. If trust falls below threshold, persuasion pauses, and the agent returns to evidence and context assessment. Influence stops feeling like pressure and starts functioning as collaboration.

Step 1: Improve Information Reliability (earn trust)

Harden the evidence tier with inspectable provenance, retrieval over a curated corpus, and honest calibration in place of confident guesswork. When sources disagree, defer to the system of record and show your work. The result is fewer hallucinations, clearer uncertainty, and answers that withstand scrutiny. If this gate does not clear, do not persuade; return to the evidence.

Step 2: Influence Through Style Alignment (earn a hearing)

Advice is accepted when it fits how people actually think. Our patented **myQ®** framework models durable differences in how individuals weigh risk, rules, and rewards, then maps those traits to reply style: tone, framing, evidence density, autonomy level, and challenge intensity. Same facts, different on-ramps. A rules-oriented user sees policy cites and checkpoints; a high-reward user gets the payoff and a clean path; a low-risk user sees limits, preview, and undo. This is personality, not persona theater—measurable, ethical, and programmable.

Step 3: Convert Persuasion to Intention (earn accountable action)

Once trust and fit are established, the agent helps users commit to a concrete plan matched to reversibility and stakes. This improves the likelihood of voluntary behavior change. Still at high-risk moments add friction through counterarguments, dual-source verification, and time-boxed holds. Low-risk, reversible steps move faster. Everything is logged in human-readable form, so accountability is visible to the user and auditable. Basically, autonomy is earned, never assumed.

Why Act Now?

As organizations scale agentic AI, users already sense which systems are generic talkers and which are dependable collaborative partners. This trust gap will widen, and trust is the entry ticket to influence. Teams that treat trust as the prerequisite, personality fit as the amplifier, and intention as the conversion step will set the new standard. They will

improve compliance, reduce risk, and lift human outcomes when it matters most, creating brand and economic advantages for those who lean in. If you are piloting or deploying AI and want to translate this method into your domain, let's talk. The clock is ticking, and the window to set the bar higher is open now.

First Up in the Series:

*Before AI can persuade, it must prove it deserves to be heard. This opening article reveals why mistrust in chatbots is rational and how real trust begins with verified evidence, calibrated confidence, and psychological fit, not flattery. You'll get a first look at our patented **myQ®** framework, which links human decision styles with AI reliability. If you want AI that earns confidence instead of demanding it, this is where the blueprint begins.*

Author Bio:

Dr. James L. Norrie is a professor of Law and Cybersecurity, and Founding Dean of the Graham School of Business, at York College of Pennsylvania (<http://www.ycp.edu>). He is the author of *Beyond the Code: AI's Promise, Peril and Possibility for Humanity* (Kendall, Hunt 2025). Learn more about our free community of interest in ethical AI at: www.techellect.com or visit www.cyberconIQ.com to learn more about AI tools to keep your employees safer online. To purchase his book, click on the QR code, or visit: <https://he.kendallhunt.com/product/beyond-code-ais-promise-peril-and-possibility-humanity>

