

# Stop Persona and Start Personality.

## 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

Article III

We do not trust cardboard cutouts; we trust people. If agentic AI expects a hearing, it must stop performing as a generic persona and align to real human personalities. Not sycophantic flattery and all-too-easy stereotypes. Rather, a disciplined fit with how individuals actually think, argue, and decide things for themselves even when collaborating with machines.

In my research for *Beyond the Code*, I asked how AI might amplify human judgment rather than replace it. That became a central thesis of the book. Now refined into the “Triple III Model,” the work shows in practice how curated information earns trust, how influence increases when aligned to personality theory, and how that, in turn, raises intention to convert sound guidance into safeguarded action. It changes decision behavior. This article focuses on that second step—**influence**—and the single insight that unlocks it: deeper personalization increases influence only when it is grounded in measurable human personality traits, not in superficial machine personas.

## Why Personas Fail & Personality Works

Personality is not a mood or a marketing segment. It is a durable, measurable, predictive pattern of traits that shows up in how people weigh **risk**, respect **rules**, and respond to **rewards**. That triad, modeled through the **myQ®** framework, explains why some of us demand authoritative policy cites and checkpoints before acting, why others pursue a clear payoff path that balances outsized risk for outsized reward, and why still others insist on a visible “undo” to test a choice’s intuitive feel before committing. Advice that meets those tendencies will be heard; advice that pushes against this intrinsic human architecture will be ignored, no matter how clever the prose. Replicating this complexity in an AI platform is possible, but it requires a deep, working grasp of both psychology and AI technologies.

By contrast, the slick personas that dominate today’s agentic AI experiments are efficient fictions at best and superficial sycophancy at worst. They compress people into tidy labels that may suit a campaign but fail at moments of consequence. A persona can pick a color palette; it cannot support a chemotherapy decision, approve a wire transfer with accountability, or help a teacher intervene at the right moment. That requires deeper alignment to personality, which the **myQ®** framework provides by giving AI agents a scientifically grounded map to align with user style.

## From Theory to Practice: Programming Personality without the Gimmicks

Patents are scrutinized for novelty and substance; they are not awarded for "vibe dials." Dismissing the **myQ®** framework as cosmetic entirely misses the point. Instead, focus on its two proven pillars of psychology theory:

**Trait-based personality science.** Decades of research show that stable traits shape how people process information, evaluate risk, and follow through. That yields a credible, durable map of meaningful interactional differences measurable across time, content and context.

**Cognitive bias and decision psychology.** The same levers that can protect us—urgency, authority, social proof, scarcity—are often exploited to hack judgment. We codify those levers, use them transparently, and hold the model to account when deploying it.

For machines to approximate human personality usefully, alignment must be programmable, testable, and improvable in the digital wild. In the Ill Model, information still comes first; fit must never outrun truth or trust collapses. Once the trust gate clears, the next gate is influence—the agent's ability to speak so its human collaborator can actually hear it, empathetically and effectively. Here is the practical training sequence:

1. **Consent or a clear cold start.** The user opts into profiling via a validated **myQ®** assessment linked to their AI profile, or the agent begins from a clearly labeled cold start, using sparse, provisional signals that improve with use.
2. **Map personality to a reply profile.** The user's **myQ®** vector across risk, rules, and rewards sets tone, framing, evidence density, autonomy level, and challenge intensity—reframing basic information into responses the user experiences as more influential.
3. **Keep trust ahead of fit.** If the information tier drops below threshold during an interaction, persuasion pauses. The agent returns to evidence, validates sources, and states uncertainty plainly.
4. **Validate fit as a living hypothesis.** The goal is improved comprehension and follow-through for this person in this context. When alignment doesn't help, the agent adjusts and re-tests.

What follows is a cycle that closes the loop with the decision-maker, now feeling more helpful to the user rather than potentially intrusive or pushy. Instead of locking people into static labels, the agent watches how alignment performs over time and adapts accordingly.

Can you imagine the difference? A generic agent talks **at** you; a personality-attuned agent works **with** you. Plans hold without late reversals because guidance fits how you weigh risk and stays within your tolerances. When a counter-argument surfaces in a

high-stakes moment, the agent slows the decision just enough to reveal the safer path. **Confidence rises for the right reasons—clear sources, reversibility, and preview-and-undo—rather than because the prose sounds certain.** And when the system errs and owns it, trust recovers.

These are not vanity signals or easy agreement. They are the recurring feedback that steers your reply profile toward your stable pattern, creating a self-training, continuously improving fit anchored to evidence, not applause. We keep one plain test in view: **WWHD—What Would a Human Do?** Personality alignment should help people act wisely, not merely agree more often. If a reasonable human would reject the nudge, so should the agent. That builds and keeps trust!

## Your Monday Morning Moves

Retire the persona approach in applications intended for human influence or decision support. Instead, stand up a personality-aware pilots with our **myQ®** framework, using opt-in and a visible transparency with users to improve trust and influence. Wire in counterarguments for high-stakes moments and track what matters: comprehension, safe-path selection, adherence without late overrides, and trust recovery after early errors. Share those results with users and market the advantages of style-aligned personalized agents versus generic AI. When people can see you earning their trust—and see the system adjusting to *them*—they will lean into its value. And you will reap the business value you always new AI could have but hasn’t yet achieved.

## Summary Conclusion

Think about the moments when judgment wobbles; a late night, a crowded inbox, a decision that matters requiring more than the clock allows. What can help in those moments? An AI-enabled tool that knows not only the relevant facts but inherent knowledge of the person weighing them. This approach truly demonstrates human care and concern expressed in collaborative AI tools.

Personality is the map; but better human decision-making is the journey. When AI honors both, particularly matching the way you balance risk, rules, and rewards to what is actually true, its counsel stops feeling like pressure and starts reading as useful partnership. That is not a tone trick or a superficial friendly veneer. It is the quiet architecture of collaboration: guidance that meets you where you are and walks with you just far enough to make the hard thing easier to do. If you are building AI, the clock is ticking; treat personality alignment as a critical best-in-class goal, and if you want help turning that principle into practice, we can show you how.

## Up Next in the Series:

The final article in the Triple Ill series explains when AI truly deserves the right to act, and when it must step back, by arguing that autonomy must be earned. You'll get a tight, practical ladder of safeguards, plus real-world rules showing how agents should ask for handoffs, offer undo paths, and always let people say "no" without penalty. If you care about AI that helps rather than hurries, this article shows the code-ready guardrails to make delegation safe, auditable, and human-centered.

### 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](http://www.techellect.com) or visit [www.cyberconIQ.com](http://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>

