

Reading Your Report

Your assessment measured two things that are usually confused for one: how well you know your domain, and how well you exercise judgment when AI can do the work. This guide helps you make sense of where you landed — and, more importantly, what to do next.

First, what this is — and isn't. The JPA is a *professional development* instrument, not a test. There are no right or wrong answers; every option you chose was defensible. It measures your *orientation* under realistic trade-offs, not your knowledge of best practice. It is a practitioner-derived diagnostic on a published validation roadmap — *not* a psychometrically validated test, and it is not designed for hiring, selection, or performance ranking. Use it to see yourself more clearly, not to label anyone.

Your two scores

Your result rests on two independent scores, each out of 5:

Domain Competency (the horizontal axis)

How deep your professional expertise runs — for the Business Analysis edition, this maps to BABOK® proficiency. It answers: *can you do the work?*

Judgment Premium (the vertical axis)

How well you exercise judgment when AI produces the work product — across the four disciplines. It answers: *can you tell when the work is right, and own the call?*

These are measured **independently** — a high score on one tells you nothing about the other. That independence is the whole point: it's possible to be highly skilled and highly over-reliant at the same time. A score of **3.0 or above** counts as "high" on each axis, and the combination places you in one of four quadrants.

The four quadrants — where you landed

LOW DOMAIN · HIGH JUDGMENT Orchestrator "Directing with wisdom" Strong judgment about when to trust AI and when to question it — domain depth still growing. Your judgment gives you leverage pure technicians lack. Do next: deepen domain knowledge strategically; substance still matters.	HIGH DOMAIN · HIGH JUDGMENT Weaver "The irreplaceable professional" Deep expertise <i>and</i> strong AI-era judgment. AI collapses the execution layer; you govern what remains. Your real output is the judgment call. Do next: lead and mentor; build the next generation of Weavers.
LOW DOMAIN · LOW JUDGMENT Interpreter "Building the foundation" Developing both axes at once — building domain literacy and AI-era judgment together, not sequentially. Do next: invest in foundational domain knowledge while experimenting with AI in low-risk contexts.	HIGH DOMAIN · LOW JUDGMENT Augmenter · highest risk "Skills without the safety net" Strong skills, underdeveloped judgment around AI. The research is clear: expertise <i>exacerbates</i> over-reliance, not mitigates it. Most likely to miss AI-introduced errors, because the output confirms what you already believe. Do next: for every AI output, ask "what is wrong, missing, or quietly assumed?" and "what did the AI not have?"

If you landed as an Augmenter, read this twice. It is not a criticism of your ability — it's the opposite. It means your domain skill is strong enough to be dangerous with AI: you can produce polished, confident, wrong work faster than anyone. This is *the Augmenter Trap* — *when AI makes you faster before it makes you wiser*. It's the most common landing spot for experienced professionals, and the most improvable.

The four disciplines — where your judgment score comes from

Your Judgment Premium score is the average of how you handled four disciplines. Each one is the AI-era expression of a habit of mind (the 4C Mindset). Your weakest of these four is where development should begin — the report names it for you.

Critical Examination FROM CLARITY

Cutting through the noise to question the premise of any AI output — surfacing what is wrong, missing, or quietly assumed. *The discipline of not accepting the machine's framing.*

Evidence Discipline FROM CURIOSITY

Forcing aspiration into arithmetic — demanding the source, the baseline, and the quantified evidence behind every claim. *The discipline of "show me the number, and where it came from."*

Political Navigation FROM CONNECTION

Reading the unspoken human reality AI cannot see — who has capacity, who blocks, who decides, what history shapes the room. *The discipline of the context that isn't in the data.*

Contextual Filtering FROM COMMITMENT TO VALUE

Filtering AI outputs and external recommendations through what will actually work *here, with these people, right now.* *The discipline of "right in general" versus "right for us."*

What to do with your result

- **Start with your weakest discipline.** Not your quadrant — your lowest of the four. That single number is the most actionable thing in the report.
- **Use the coaching prompts.** Your results screen generates AI coaching prompts personalised to your profile — "Develop My Weakest Discipline" is the one to use first.
- **Re-read the scenario breakdown.** The table shows how you scored question by question. Look for the ones where you chose the impressive-sounding option — those are where the Augmenter Trap lives.
- **Retake it in a few months.** Judgment is a practice, not a trait. The instrument is designed to show movement.

About a validity notice. If your report showed a validity warning, it simply means some answers came faster than the scenarios can be read (they run 150–250 words). It doesn't invalidate you as a professional — it means *that particular run* may not reflect your real judgment. Retake it unhurried; the result will mean more.

The JPA is a practitioner-derived instrument on a defined validation roadmap. Your result is a mirror for development, not a measure of your worth, and not a basis for selection or ranking. For the evidence base and validation status, see the JPA Methodology paper.

© 2026 Susan K.L. Yu · Altschool Australia · Judgment Premium™