

Honest assessment

- [Placing a person on the levels](#)
- [Failure modes](#)
- [Validation at scale](#)

Placing a person on the levels

The method lives or dies on honest assessment — and the levels are built to make honesty achievable. They're defined by *observable behaviour*: not "how good is this person?" but "what happens when they do the work?"

The tells, level by level:

Beginner — the work gets done, but with help or documentation *along the way*, not just at the edges. Fresh from a course or a book.

Confident — most tasks delivered independently, reaching for docs only on the less common parts. The work gets done, by them, routinely.

Expert — almost everything handled without help; lookups only for obscure cases. The strongest single tell: **they can teach it**. Someone who can walk a colleague from beginner to confident in a skill is an expert in it.

(And the standing footnote: **even experts look things up**. Skills evolve; needing materials at any level is normal and expected. For small no-level skills, none of this applies — they're a plain yes/no.)

The scale is fixed; the sources vary

Because the tells are behavioural, *anyone who has seen the work* can read them — which is what lets the same scale accept assessments from many sources:

- **Manager assessment** — the default in most professional environments; managers see delivery month after month, which is exactly what the tells describe.
- **Self-assessment** — the natural default for personal use, and a valuable input everywhere; nobody knows better how often help was needed along the way.
- **Peers and 360 feedback** — colleagues who work alongside someone often have the clearest view of all, especially for skills a manager rarely observes directly.
- **Credentials and certifications** — solid evidence for beginner (the basics are demonstrably known) and useful signals beyond; they certify knowledge, so pair them with delivery for the higher levels.

- **Teaching-back** — the built-in test that needs no interpretation: teach it, and expert is demonstrated, not claimed.

Which sources carry the decision is a cultural and organisational choice — the method doesn't mandate a mix. What it insists on is only this: every source answers the *same* behavioural question against the *same* scale. That's what keeps a level meaning one thing no matter who assessed it — and it's why disagreement between sources is useful rather than awkward: it's a small, concrete conversation about one skill, settled by looking at the work.



Failure modes

Every assessment source has a bias signature. Knowing them is most of the defence — and it's why the method prefers levels that several sources can confirm.

Inflation. A level claimed (or granted) above the behaviour. It feels harmless, even kind — until the map assigns work that assumes foundations that aren't there. The cost arrives as frustration: tasks that feel impossibly hard, slipping deadlines, and a person quietly concluding they're not good at this — when the truth is only that a level got inflated and the missing foundation was never scheduled. Inflation doesn't skip the learning; it moves it to the worst possible moment.

Modesty. The same error in reverse, with a gentler face. Rate someone beginner where they're confident — whether from their own humility or an assessor playing it safe — and the map routes them around opportunities they were ready for: the stretch project, the mentoring role, the promotion case. The map can only open doors it believes someone can walk through.

Source-specific tilts. Self-assessments drift toward whichever of the two errors matches the person's temperament. Manager assessments lean on what the manager sees — skills exercised out of view get under-rated, recent wins outshine steady delivery. Credentials go stale: the certificate is forever, the skill isn't. None of this disqualifies a source; it just means no single source should be the whole story on a skill that matters.

Three habits keep placement honest, whoever is assessing:

- **Anchor on the behavioural tells.** Help along the way, or only at the edges, or almost never? That's a fact about last month's work, not a judgment of worth — and it's the same question for every source.
- **Be specific.** Broad skills invite vague ratings in any direction; "social-media posting: confident, SEO basics: beginner" leaves little room for ego *or* halo.
- **Correct cheaply, early.** The map has no memory of pride. Fixing a level costs nothing; living with a wrong one costs every plan built on it.



Validation at scale

Assessing a handful of people honestly is a discipline. Assessing fifty — while levels feed reviews, promotions, and eventually pay — is a system design problem. The pressure is predictable: **once levels touch money, levels inflate**. Not because people are dishonest, but because incentives lean on the scale — on self-ratings and manager ratings alike — person by person, a little at a time.

The counterweight is making every level *a claim others confirm*:

Peer validation. A level isn't declared by one voice — whether the person's own or their manager's; the people who work alongside them corroborate it. This isn't a tribunal — most of the time it's a colleague nodding "yes, she delivers that on her own." The point is that levels become shared observations, not private declarations. Disagreement is a conversation about one skill at one level — small, concrete, and usually resolved in minutes.

Skill champions. For each critical skill, name a champion — someone at expert level who can answer questions, mentor others, *and validate assessments* in that skill. Champions give every skill a reference point: an agreed example of what expert actually looks like here. When someone wonders whether they're confident yet, the champion has seen enough of the skill to say. (Champions carry more weight in the team-development story too — see the team chapter.)

Teaching-back as the standing bar. The confident→expert test scales beautifully, because teaching is public. A claimed expert who has walked two colleagues to confident needs no further audit.

Handled this way, validation barely feels like control. It feels like what it is: a team keeping its shared map accurate, because everyone's plans are drawn on it.

