

WHITE PAPER • 2026

Four questions. **One decision.**

Performance, Potential, Readiness and Reputation are not synonyms. Telling them apart is the difference between a people decision that holds and one that gets justified after the fact.

Four different questions. **Never synonyms.**

In most organizations the same people get described with words that sound interchangeable: she performs, she has potential, she is ready, she is well regarded. These are four different statements, measured in different ways, leading to different decisions. Confusing them is the most expensive mistake in talent management, because it moves people on the basis of the wrong question.

01 · POTENTIAL

Who can grow?

The likelihood that a person will develop, over time, the competencies required by more complex roles. It rests on foundations - cognitive ability, personality traits - and on enablers such as learning agility and motivation to grow.

→ Long-term development decisions

02 · READINESS

Who is ready now?

How many of the competencies required by a more complex role a person already holds, today. It measures the present against a future role, not the trajectory.

→ Succession and promotion decisions

03 · PERFORMANCE

Who excels in the current role?

Current competencies against the requirements of the current role. It describes the present within today's perimeter and anticipates nothing about different roles.

→ Recognition, reward, retention

04 · REPUTATION

How is the behaviour perceived?

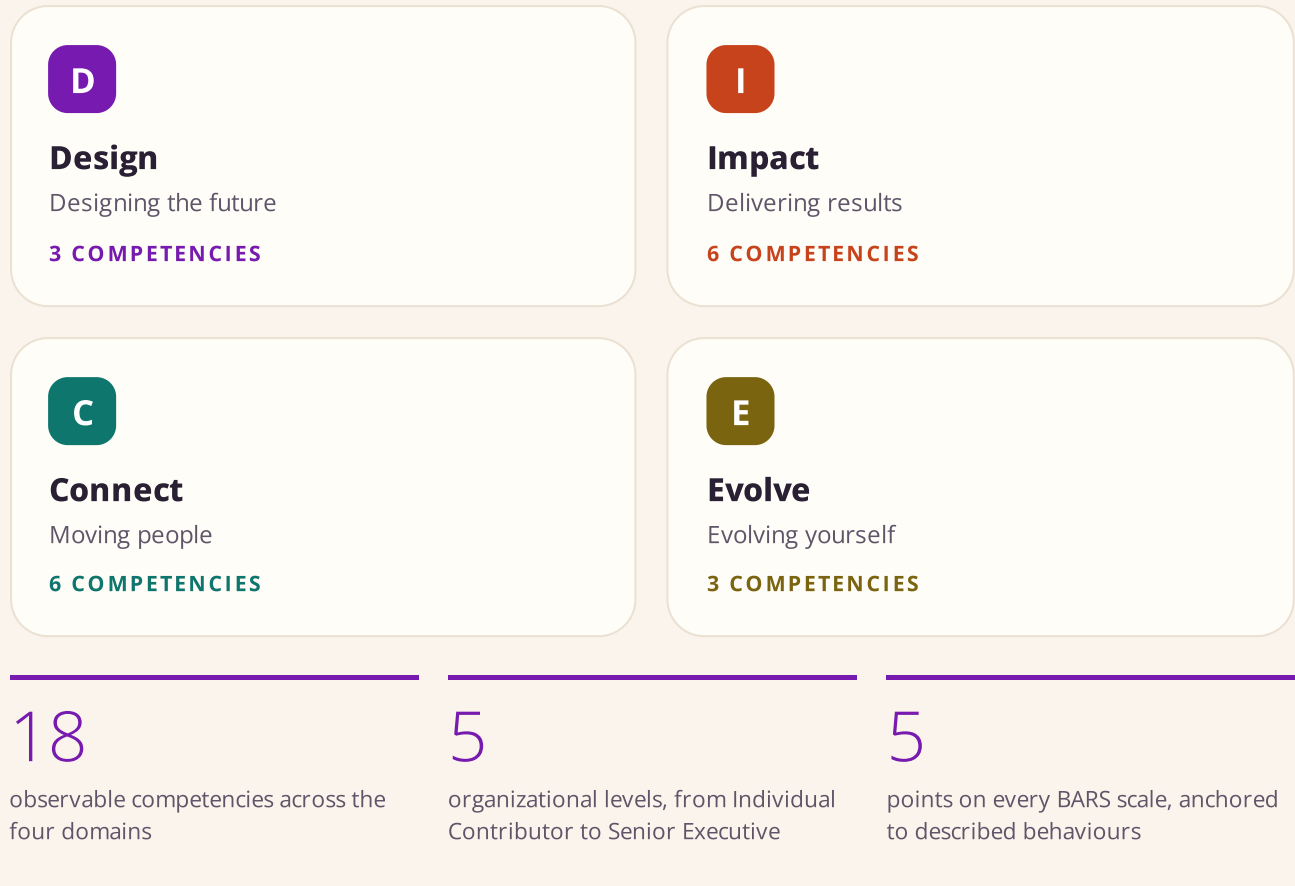
Current competencies as they are experienced and perceived by managers, peers, direct reports and other stakeholders. It is the data that explains why strong performance can go unrecognised.

→ Feedback and behavioural development

The practical test. Take the last person you promoted and ask which of the four questions you actually had data on. If the answer is "on performance", you decided about a future role by looking at the past one.

One **competency model** beneath every measure

For four different measures to stay comparable across people, roles and moments in time, they must rest on the same vocabulary. DICE is the proprietary Quint model: it translates leadership requirements into observable behaviours, specific to each organizational level and anchored to described scales.



Why anchored scales matter

BARS stands for **Behaviorally Anchored Rating Scale**: each point on the scale is not an adjective but a described behaviour. It is the difference between "good influencing skills" and a sentence that states exactly what that person does, in which situations and with what consistency. It is also what makes a score explainable to the person being assessed, and defensible to anyone who challenges the decision.

What if we already have our own model?

Potential and Readiness can also be assessed against the competency model your organization already uses. The model is the vocabulary; what does not change is the method by which evidence is collected.

AI-enabled, **not AI-only**

Technology makes it possible to measure at scale what used to be measured for a few. It does not make the person who decides redundant. When a score is generated with AI support, the person receiving it has a right to know what evidence it rests on - and the person using it needs to be able to defend it.

*Technology makes the measurement possible. **People** remain the ones who decide.*

Four questions to ask before you sign

1

Which construct is this instrument validated on? An instrument validated on performance says nothing about potential, however sophisticated it is.

2

What does the assessed person see? If feedback to the person does not exist or is not intelligible, the data will produce resistance, not development.

3

How do you explain a low score? If the answer is "the model says so", the score is defensible neither in a dispute nor in a conversation.

4

Do results stay comparable over time? Without the same model and the same scale, two measurements a year apart are two opinions.

Start from the **question**, not from the product

The first step is not choosing an instrument: it is writing down, in one sentence, the decision you face over the next twelve months. The question determines the construct, the construct determines the method. Done in that order, the scope stays small and the result is verifiable.

HO	HORIZON - Potential Assessment	"Who can grow into more complex roles?"
SW	SWITCH - Readiness Assessment	"Who is ready now for a bigger role?"
FV	FULL VIEW - Comprehensive Assessment	"Can we decide promotions better?"
CO	COMPASS - Self-awareness at Scale	"How do we build self-awareness for everyone?"
MI	MIRRORS - Real Reputation at Work	"How do our managers receive feedback?"
GL	GROWTH LAB - Turning Insight into Action	"How do we turn insight into practice?"
IM	IMPACT - Measuring Behavioral Change	"Did the development actually work?"
LI	LINKS - Mapping Collaboration Networks	"Where does collaboration really happen?"

Let's talk about your next people decision

Thirty minutes with the Quint team: bring a real question, leave with the construct to measure and a concrete path.

welcome@quint.org · +39 342 8137512 · quint.org