

WHITE PAPER • 2026

Quattro domande. **Una decisione.**

Prestazione, potenziale, prontezza e reputazione non sono sinonimi. Distinguerle è la differenza fra una decisione sulle persone che regge e una che si giustifica a posteriori.

Quattro domande diverse. **Mai sinonimi.**

Nella maggior parte delle organizzazioni le stesse persone vengono descritte con parole che sembrano intercambiabili: è brava, ha potenziale, è pronta, è ben vista. Sono quattro affermazioni diverse, misurate in modi diversi, che portano a decisioni diverse. Confonderle è l'errore più costoso del talent management, perché sposta le persone sulla base della domanda sbagliata.

01 · POTENZIALE

Chi può crescere?

La probabilità che una persona sviluppi nel tempo le competenze richieste da ruoli più complessi. Si fonda su fondamenta - capacità cognitive, tratti di personalità - e su abilitatori come learning agility e motivazione alla crescita.

→ **Decisioni di sviluppo a lungo termine**

02 · PRONTEZZA

Chi è pronto ora?

Quante delle competenze richieste da un ruolo più complesso una persona possiede già, adesso. È la misura del presente rispetto a un ruolo futuro, non della traiettoria.

→ **Decisioni di successione e promozione**

03 · PRESTAZIONE

Chi eccelle nel ruolo attuale?

Le competenze attuali rispetto ai requisiti del ruolo attuale. Racconta il presente nel perimetro di oggi e non anticipa nulla su ruoli diversi.

→ **Riconoscimento, retribuzione, retention**

04 · REPUTAZIONE

Come viene percepito il comportamento?

Le competenze attuali così come vengono vissute e percepite da manager, colleghi, collaboratori e altri stakeholder. È il dato che spiega perché una prestazione buona può non essere riconosciuta.

→ **Feedback e sviluppo dei comportamenti**

La prova pratica. Prendete l'ultima persona promossa e chiedetevi su quale delle quattro domande avete davvero un dato. Se la risposta è "sulla prestazione", avete deciso su un ruolo futuro guardando il ruolo passato.

Un solo **modello di competenze** sotto ogni misura

Perché quattro misure diverse siano confrontabili fra persone, ruoli e momenti, devono poggiare sullo stesso vocabolario. DICE è il modello proprietario Quint: traduce i requisiti della leadership in comportamenti osservabili, specifici per livello organizzativo e ancorati a scale descritte.

D

Design

Progettare il futuro

3 COMPETENZE

I

Impact

Produrre risultati

6 COMPETENZE

C

Connect

Muovere le persone

6 COMPETENZE

E

Evolve

Evolvere se stessi

3 COMPETENZE

18

competenze osservabili, distribuite sui quattro domini

5

livelli organizzativi, da Individual Contributor a Senior Executive

5

punti in ogni scala BARS, ancorati a comportamenti descritti

Perché le scale ancorate contano

BARS significa **Behaviorally Anchored Rating Scale**: ogni punto della scala non è un aggettivo ma un comportamento descritto. È la differenza fra "buona capacità di influenza" e una frase che dice esattamente che cosa quella persona fa, in quali situazioni e con quale continuità. È anche ciò che rende un punteggio spiegabile davanti a chi viene valutato, e difendibile davanti a chi contesta la decisione.

E se abbiamo già un nostro modello?

Potenziale e prontezza possono essere valutati anche sul modello di competenze già adottato dalla vostra organizzazione. Il modello è il vocabolario: quello che non cambia è il metodo con cui si raccoglie l'evidenza.

AI-enabled, **non AI-only**

La tecnologia rende possibile misurare su larga scala ciò che prima si misurava solo su pochi. Non rende superflua la persona che decide. Quando un punteggio è generato con supporto AI, chi lo riceve ha diritto di sapere su quale evidenza poggia - e chi lo usa ha bisogno di poterlo difendere.

*La tecnologia rende possibile la misura. **Le persone** restano quelle che decidono.*

Quattro domande da fare prima di firmare

1

Su quale costruito è validato questo strumento? Uno strumento validato sulla prestazione non dice nulla sul potenziale, per quanto sofisticato sia.

2

Che cosa vede chi viene valutato? Se la restituzione alla persona non esiste o non è comprensibile, il dato produrrà resistenza, non sviluppo.

3

Come si spiega un punteggio basso? Se la risposta è "lo dice il modello", il punteggio non è difendibile in un contenzioso né in un colloquio.

4

I risultati restano confrontabili nel tempo? Senza uno stesso modello e una stessa scala, due misurazioni a un anno di distanza sono due opinioni.

Partite dalla **domanda**, non dal prodotto

Il primo passo non è scegliere uno strumento: è scrivere in una frase la decisione che avete davanti nei prossimi dodici mesi. La domanda determina il costrutto, il costrutto determina il metodo. Fatto in quest'ordine, il perimetro resta piccolo e il risultato è verificabile.

HO	HORIZON - Potential Assessment	"Chi può crescere verso ruoli più complessi?"
SW	SWITCH - Readiness Assessment	"Chi è pronto ora per un ruolo superiore?"
FV	FULL VIEW - Comprehensive Assessment	"Possiamo decidere meglio le promozioni?"
CO	COMPASS - Self-awareness at Scale	"Come sviluppiamo l'auto-consapevolezza per tutti?"
MI	MIRRORS - Real Reputation at Work	"Come ricevono feedback i nostri manager?"
GL	GROWTH LAB - Turning Insight into Action	"Come trasformiamo l'insight in allenamento?"
IM	IMPACT - Measuring Behavioral Change	"Lo sviluppo ha davvero funzionato?"
LI	LINKS - Mapping Collaboration Networks	"Dove avviene davvero la collaborazione?"

Parliamo della vostra prossima decisione sulle persone

Trenta minuti con il team Quint: portate una domanda reale, uscite con il costrutto da misurare e un percorso concreto.

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