



# Is Your Field Force Just Certified or Field-ready?

Rubrics That Predict Field Readiness

OCTANE

Every field trainer has seen it: a rep aces the certification exam, hits every knowledge checkpoint, and then freezes the first time a healthcare professional (HCP) pushes back on a message or data point. Being able to deliver that knowledge to the right customer at the right time is distinct from demonstrating knowledge of content, and most training programs still measure the first while assuming it predicts the second.

Certification and assessment scores tell you what a rep can recall in a controlled environment. They don't tell you what a rep can do in a seven-minute call with a skeptical HCP who wants the confidence interval, not the headline. That gap is where field readiness lives, and it's exactly the gap a well-built verbalization rubric is supposed to close.

When rubrics are vague or inconsistent, the costs show up fast. Most important are longer ramp times, which already average more than 11 months in pharma to get reps to the average level of performance.<sup>1</sup> Other issues can include managers scoring reps inconsistently or coaching conversations that amount to "good job" or "needs work" with nothing actionable underneath. We've seen it firsthand: managers checking off "delivers efficacy results sufficiently" because the rep read it straight off the visual aid, with the rubric having no way to register how that landed in the actual conversation.

## What "Good" Looks Like

A strong verbalization rubric shares a few traits regardless of therapeutic area or company:



### It's behaviorally anchored, not a generic scale.

Most rubrics still default to a 1–5 Likert scale with labels like "communication skills" or "product knowledge." The problem is that a 3 means something different to every evaluator. Behaviorally Anchored Rating Scales (BARS) solve this by defining what each score point looks like in observable terms.<sup>2</sup>



### It's specific enough to observe, not interpret.

"States the primary endpoint and confidence interval accurately" is something two evaluators can agree on. "Understands the data" is not. If a criterion requires the evaluator to infer understanding rather than observe behavior, it's not ready to score against.



### It's calibrated.

Even a well-written rubric falls apart if trainers and managers apply it differently. Inter-rater reliability isn't a nice-to-have — it's the thing that determines whether a score means anything at all.<sup>3</sup>



### It's weighted to reflect what really matters.

Not every part of the rubric deserves equal weight. While delivery is still a vital part of the rep's role, compliance and scientific accuracy should carry more weight than polish or persuasive delivery, even though delivery is often what's easiest to score.



### It has real tiers.

"Needs improvement / meets expectations / exceeds expectations" is only useful if each tier is defined by different behavior, not just a different adjective.

## What Everyone Gets Wrong

*Everyone builds the rubric to score the call. The real job is to build a rubric that predicts the next one.*

Most rubrics are backward-looking by design: built to produce a defensible, consistent score for the performance that just happened. Calibration and inter-rater reliability matter but consistency isn't prediction. A rubric that only scores well is answering the wrong question. The question that really matters isn't "how did this rep do?" It's "does this score tell us anything about how this rep will do in front of an HCP who wasn't in the script?"

That's the gap almost nobody closes, because closing it means treating the rubric as a predictive instrument, not a grading instrument. In practice, that requires three things most organizations skip:

- Validating the rubric against actual field outcomes (e.g., ramp time, call effectiveness, HCP-reported outcomes), not just checking it for internal consistency.
- Revisiting the rubric when it stops predicting. Most organizations build a rubric once, calibrate it once, and never ask whether the reps who score well are the reps who are succeeding in the field 18 months later.
- Building the rubric around the moments that clearly separate reps in the field (e.g., unscripted pushback, off-label questions, objection handling under time pressure) rather than focusing too much around the moments that are easiest to observe and score in a controlled role-play.



This is the piece other frameworks skip, because it's the hardest to operationalize. It's much easier to write "understands the data" than it is to build the feedback loop that tells you, six months later, whether that criterion was the one that mattered. "Understands the data" should really be "can reprioritize which data point to lead with when the HCP's objection shifts mid-call." The first is easy to score. The second is the thing that separates reps in the field and the thing that almost no rubric is built to capture. That loop of rubric to field data back to rubric is chronically missing, and closing it is what turns a scoring tool into an actual predictive instrument.

## Anatomy of a Verbalization Rubric

The strongest rubrics we've seen tend to organize around five domains:

### 1. Scientific accuracy and on-label adherence

This is the domain everyone gets right at a surface level and wrong at a granular one. Most rubrics stop at “understands the data,” but accuracy under a controlled recitation and accuracy under interruption are different skills, and only one of them predicts field performance.



#### What needs to be observed:

- Does the rep state the primary endpoint, the comparator, and the confidence interval unprompted, not just when asked?
- When challenged (“that’s not what I read”), does the rep hold the accurate data point, or does the data quietly soften to match the pushback?
- Does the rep stay on-label without visibly reaching for a script, and does the on-label language sound like fluency, or like a guardrail they’re nervous about tripping over?

The failure mode here isn’t reps who don’t know the data. It’s reps who know it cold in isolation and lose precision the moment a real question knocks them off the rehearsed sequence.

We’ve seen this play out with a number of clients whose rubric consistently scored reps as “meets expectations” on competitive positioning. But field rides three months later told a different story: the same reps could recite the differentiators perfectly in a roleplay but froze or fell back on generic talking points when an HCP brought up a competitor’s data. The rubric was measuring whether a rep could list the differentiation points, not whether they could deploy the right one in real time against a live objection.

### 2. Message flow

This domain is usually scored as “did they cover the key messages in order,” which measures memorization, not flow. Real message flow is whether the rep can move through the core narrative while responding to what the HCP actually says — not deliver the narrative and handle interruptions as a separate skill.



#### What needs to be observed:

- Can the rep return to the core message after a detour, without it sounding like they’re rebooting a script?
- Does the sequencing still make clinical sense when a question forces the rep to cover point three before point one?
- Is there a visible “reset” (e.g., a verbal stumble or restart) when the conversation goes off the expected path?

The tell here is usually in the transitions, not the content. A rep with strong message flow makes a redirected conversation sound like it was always going to go that way.

### 3. Adaptive communication

This is the domain most rubrics gesture at (“communicates effectively with different HCP styles”) without ever defining it in observable terms — the same “understands the data” problem named earlier, just relocated to tone instead of content.



#### What needs to be observed:

- Does the rep change what they lead with based on signals from the HCP (e.g., a data-driven physician gets the confidence interval up front, a time-pressured one gets the bottom line first) or do they run the same sequence regardless of who's in front of them?
- Does the rep respond differently to skepticism versus genuine curiosity, or do they handle both with the same rehearsed reassurance?
- When the HCP gives almost no signal — flat affect, minimal engagement — does the rep have a default mode that's still effective, or do they default to over-talking?

Some rubrics solve this by anchoring the evaluation to a short scenario up front (e.g., a specific HCP profile and objection the rep is responding to) rather than leaving “adaptive” to be scored against whatever generic physician the evaluator happens to be picturing. A few clients have taken this further, building out a small set of scenarios and having the rater select the one that best matches the individual rep's actual call before scoring begins. That way, adaptive communication, message flow, and presence are all being scored against the same fixed reference point the rep was responding to, not an abstract standard.

### 4. Compliance and regulatory guardrails

Usually scored as a pass/fail checklist — did the rep say anything off-label or did they use an approved visual aid. That's necessary but insufficient. The higher-value observation is whether compliance holds up under conversational pressure, not just under normal delivery.



#### What needs to be observed:

- Is confidence tied to accurate content, or is it present regardless of what's being said? This criterion should be cross-checked against Domain 1 (scientific accuracy and on-label adherence) — presence should never be scored in isolation from accuracy.
- Is the redirect back to on-label territory smooth, or does it sound like a compliance flag going off mid-conversation?
- Does the rep know the difference between declining to answer and redirecting to an appropriate resource — one preserves the relationship, the other can read as evasive?

This is the domain that is most likely to be scored generously, because a rep who “didn't say anything wrong” often gets full marks even when they handled the moment clumsily.

## 5. Confidence or executive presence

This is the domain most responsible for the “sounds sure, is wrong” failure mode named earlier in this piece — which is exactly why it needs the tightest behavioral anchors, or it becomes the domain that rewards the wrong reps.



### What needs to be observed:

- Is confidence tied to accurate content, or is it present regardless of what's being said? This criterion should be cross-checked against Domain 1 — presence should never be scored in isolation from accuracy.
- Does the rep recover from being caught without an answer gracefully (acknowledging and following up), or cover with false certainty?
- Is presence consistent across a friendly HCP and a hostile one, or does it collapse under real pushback?

Scored on its own, this domain is the one most likely to produce a confident rep who's wrong. Scored against Domain 1, it becomes one of the more useful predictors of how a rep will perform when a call doesn't go their way.

## Common Pitfalls

An observer may be subjectively rewarding confidence over accuracy when a rep who sounds sure of themselves in their delivery ends up scoring well on a rubric (even when the content is off). Studies have found that when participants feel highly confident about their responses, they are correct only 75% of the time.<sup>4</sup> Confidence in response does not always equate to accuracy, so a coach needs to keep that in mind when assessing a rep through a rubric.

Overloaded rubrics are another problem; the usual 10 or 15 criteria are more than any evaluator can hold in their head while actively listening. On the other hand, underloaded rubrics may miss the nuance that separates a good verbalization from a great one. Our understanding of cognitive load suggests working memory can hold only 5–9 items at any one time, which may be an ideal target for developing rubrics.<sup>5</sup>

Compound that with a missing calibration process and you get managers scoring the same performance in different ways. Lastly, rubrics that get built once and never checked against actual field outcomes tend to keep measuring the wrong things long after they've stopped being useful.



## Building and Validating a Rubric

Good rubric design starts before anyone writes a single criterion. **Start by naming the actual problem:** what new development or performance gap prompted the need for this rubric in the first place?

From there, build it with the right people in the room: learning and development, training managers, stakeholders, and ideally someone close to top-performer behavior in the field.

Pilot the draft with a small group and check for inter-rater agreement before rolling it out broadly. If it's going to scale across a large field force, plan for training the trainers on how to use it.

Perhaps most critically, feed it back into coaching. A rubric that only produces a score has failed; a good one gives managers language they can use in a coaching conversation.

Rubrics also need to flex with the moment, not just with time. A new messaging campaign, a competitive shift, or a rep moving into a new therapeutic area all change what "good" verbalization looks like, even if the underlying rubric structure hasn't changed. Some clients keep a stable core (the five domains, the weighting) and swap scenario-specific criteria or benchmarks tied to whatever's live in the field right now, rather than treating the rubric as one static instrument for every context.

This is the process we run with clients: a stakeholder session to identify/name the actual field gap, a small pilot to check for agreement before scaling, and a built-in measurement point (e.g., 6 months) to check the rubric against real data rather than assuming it's still working.

## Where AI and Technology Fit

AI-assisted role-play scoring is increasingly part of this picture, and it's a legitimate tool for consistency at scale. It can apply the same rubric the same way, every time, across hundreds of reps, which is a real advantage over human raters, whose consistency degrades with fatigue and caseload. But consistency at scale is not the same claim as predictive validity, which is the point above.

Technology reduces variability — it doesn't decide what "good" looks like. Where it adds real value is in the feedback loop: using speed-to-proficiency data to see whether the rubric is predicting field performance and using that data to adjust accordingly.

## From Score to Standard

A good verbalization rubric does more than score a rep's performance. It defines what "good" actually looks like, clearly enough that a manager, a peer, and an AI-assisted platform would all reach the same conclusion watching the same conversation. That only happens when the rubric is behaviorally anchored, weighted toward what genuinely predicts field success, and calibrated across the people using it. Technology can help extend that consistency to scale, but it can't replace the upfront work of deciding what matters and defining it in terms people can observe.

The organizations that get this right treat the rubric as something living, not a document built once and filed away, and they check it against real field outcomes as products, competitors, and HCP expectations shift. Get the rubric right, and everything downstream — coaching, certification, ramp time — gets a lot easier to get right as well.

If you're not sure whether your current rubric is still predicting field performance, or measuring the moments that actually matter, that's usually a quick thing to check. We're happy to help you assess your training ecosystem and walk through what that would look like for your program.

## References

1. The Top Priority for Sales Leaders: Improving Sales Enablement and Training in Pharma. PharmExec. Published June 2026. Accessed July 16, 2026. <https://www.pharmexec.com/view/priority-sales-leaders-improving-enablement-training>
2. Behaviorally Anchored Rating Scales (BARS) – Definition, Structure & Practice. Aivy. Accessed July 13, 2026. <https://www.aivy.app/en/lexicon/behaviorally-anchored-rating-scales-bars>
3. Human Feedback Quality: Rater Training and Inter-Rater Reliability. ACCESSIBLE. Accessed July 13, 2026. <https://accessible-eu.org/human-feedback-quality-rater-training-and-inter-rater-reliability>
4. Overconfidence. iResearchNet. Accessed July 16, 2026. <https://psychology.iresearchnet.com/social-psychology/decision-making/overconfidence/>
5. What is cognitive load? The eLearning Coach. Updated November 11, 2024. Accessed July 16, 2026. <https://theelearningcoach.com/learning/what-is-cognitive-load/>