



SAMPLE REPORT · ILLUSTRATIVE COMPANY

The Revenue Execution Audit

A full, unedited example of the diagnostic every xDR Coach engagement begins with: real methodology, real maths, and a fictional company built to show the format.

PREPARED FOR

Meridian Cloud Solutions

AUDIT PERIOD

Q2, 6 Sellers, 60 Audited Calls

LEVEL ASSESSED

Level 1, Foundational Execution

Revenue is healthy. Repeatability is not.

The story in one page, before any table asks you to interpret it.

Meridian's sales team closed the quarter at **\$283,200** against a **\$278,764** target: **101.6% attainment**. On the surface, this is a team hitting its number.

Underneath that number, the picture changes. Across the six sellers in this sample, average execution capability sits at **70.3 out of 100**, below the 75-point threshold this methodology treats as reliable, repeatable execution.

- **64.7% of quarterly revenue** was produced by sellers below that threshold. The majority of this result is being carried by capability that has not yet been demonstrated as consistent.
- **Connor Blake** produced the team's strongest individual result, 116% attainment, while showing one of the weakest execution scores in the sample.
- **Priya Anand** shows both the strongest execution score and the strongest attainment: direct proof, from inside this team, that the standard is achievable.

This is not a revenue performance problem. It is a revenue repeatability problem. The team is hitting target on the strength of a small number of sellers, not a demonstrated, sustained standard of execution.

If nothing else on this page is read, this is the part that matters.

01 WHAT IS HAPPENING? Revenue performance is healthy: 101.6% attainment. Execution consistency is not. 64.7% of that revenue is produced below the execution standard.	02 WHY DOES IT MATTER? The team is currently relying on individual performance, not a consistent operating standard. The strongest single result this quarter comes from one of the two weakest execution scores in the sample.	03 WHAT SHOULD LEADERSHIP DO? Do not train the whole team. Fix two recurring behaviours in two specific sellers first. Investigate one high performer before coaching him, then re-measure.
---	--	--

Revenue versus Target

METRIC	VALUE
Quarterly target	\$278,764
Quarterly actual	\$283,200
Attainment	101.6%
Sellers in sample	6
Calls audited	60 (10 per seller)

Read alone, this supports a positive narrative. Section 4 places it beside execution capability: the evidence that determines whether it is repeatable.

Where each seller actually sits

Plotting commercial result against execution capability separates four distinct situations, and reveals the one that revenue dashboards alone cannot show.



Priya Anand

Priya records the highest execution score in the sample **(91)** alongside the highest attainment **(109%)**: direct, internal proof that the standard this audit measures against is achievable inside this specific team, without appeal to territory, tenure, or deal luck.

Her full ten-call audit identifies precisely which protocols she demonstrates consistently, and whether those behaviours are teachable to the rest of the team or substantially tied to individual factors that do not transfer. That determination is the first analysis run in a live engagement.

Recommended use: Priya's calls become the internal reference recording for coaching the rest of the team. Proof of "what this looks like when it works here" is more persuasive than any external example.

Two patterns, evidenced across the sample

01. Discovery Talk-Time

Observed in 5 of 6 sellers' audited samples. Average seller talk time across flagged calls measured 74%, well above the 65% threshold this methodology treats as the point a rep stops eliciting and starts talking over the buyer's own diagnosis.

Evidence (Elena Torres, Call 6): "You said your current process is 'a bit clunky', and I moved straight into our workflow builder instead of asking what clunky actually costs your team each week."

02. Permission-Seeking Language

Observed in 4 of 6 sellers' audited samples. Rising-inflection phrasing at agenda-setting and pricing moments ("does that make sense?", "is that okay?") recurs at moments that should signal authority instead.

Evidence (Marcus Webb, Call 4): "We'll cover onboarding and pricing today. Is that alright with you?"

One card each. Full evidence lives in each seller's individual Lab Report.

Priya Anand

\$58,200 rev • 109.0% attain • 91 exec

INTERNAL BENCHMARK

Highest execution and attainment in the sample. Reference standard for team coaching.

Connor Blake

\$61,400 rev • 116.0% attain • 68 exec

INVESTIGATE FIRST

Largest revenue and largest modelled exposure in the sample. Causality between execution and result not yet established.

Elena Torres

\$39,600 rev • 91.0% attain • 55 exec

PRIORITY 1

Weakest execution score in the sample. Clearest, best-evidenced intervention case.

Marcus Webb

\$37,900 rev • 89.0% attain • 61 exec

PRIORITY 2

Second-clearest intervention case, sharing the same primary constraint as Elena.

Sasha Kim

\$41,800 rev • 104.0% attain • 76 exec

WATCH

Just above the pass mark. Monitor for sustained consistency before next review.

Ryan Souza

\$44,300 rev • 96.0% attain • 71 exec

WATCH

Below pass mark but not top priority this cycle. Re-assess at next audit.

Three figures. Only one depends on an assumption.

A CRO can reasonably ask why a modelled percentage should be trusted. The honest answer is that it should not carry the whole argument, so this model does not ask it to.

A. Execution Exposure (observed, not modelled)

How much current revenue is produced by sellers below the 75-point pass mark. No assumption required: this is the revenue and score data already collected.

\$183,200 of quarterly revenue (64.7%) was produced by sellers scoring below 75: Elena, Marcus, Connor, Ryan, and Sasha.

B. Modelled Opportunity (the one modelled figure)

This runs on the same 0–100 execution scale every call in this audit is already graded against. 75 is the professional operating benchmark. 100 is complete execution of the standard. The assumption, that each point of shortfall below 100 corresponds to roughly 0.5% of that seller's revenue in unrealised opportunity, is deliberately conservative and bounded: it can never exceed half of any seller's actual revenue, at any score.

SELLER	REVENUE	SCORE	GAP	CONSERVATIVE	FULL BENCHMARK
Elena Torres	\$39,600	55	45	\$4,455	\$8,910
Marcus Webb	\$37,900	61	39	\$3,695	\$7,390
Connor Blake	\$61,400	68	32	\$4,912	\$9,824
Ryan Souza	\$44,300	71	29	\$3,212	\$6,424
Sasha Kim	\$41,800	76	24	\$2,508	\$5,016
Priya Anand	\$58,200	91	9	\$1,310	\$2,619

Team total, monthly: **\$20,092 conservative** and **\$40,183 at full benchmark closure**. Annualised: approximately **\$241,000** and **\$482,000**. Neither figure is a revenue guarantee.

This is not spread evenly. **65% of it** (\$13,062 of \$20,092 conservative, monthly) sits with the three lowest-scoring sellers, and traces back to two recurring capability gaps, not six unrelated problems.

C. Revenue at Risk (a risk metric, not an opportunity)

How much current revenue depends on a seller whose result is not yet supported by demonstrated, repeatable execution: the profile most exposed if that seller leaves, is promoted, or market conditions shift.

\$61,400 in quarterly revenue (Connor Blake) is currently produced by the sample's clearest high-performance, weak-execution profile: a commercial outlier requiring investigation, not yet a confirmed cause-and-effect finding.

Ranked by exposure, frequency, and evidence strength, not dollar value alone

PRI.	SELLER	EXPOSURE/M0	PRIMARY CONSTRAINT	INTERVENTION	RE-TEST
1	Elena Torres	\$4,455	Discovery / talk-time	Intensive coaching	Day 60
2	Marcus Webb	\$3,695	Discovery / talk-time	Intensive coaching	Day 60
3	Team-wide	n/a	Permission-seeking language	Systemic call standard	Day 90
4	Connor Blake	\$4,912	Unconfirmed	Diagnostic review, not yet training	Day 30 (investigation)

Connor Blake carries the single largest modelled figure in the sample and ranks fourth, not first. His revenue result is the team's best; his execution score is genuinely weak. Coaching him on the wrong assumption risks disrupting a result that may not be behaviour-driven at all. Investigation is the correct first step, not training.

If leadership can only fund two coaching interventions this quarter, rows 1 and 2 are those two. Row 4 is not a fourth intervention: it is the investigation that should happen before anyone spends coaching hours on Connor specifically.

Restraint is part of the recommendation

- **Do not retrain the entire team.** 65% of modelled opportunity concentrates in three sellers and two behaviours. A team-wide programme spends most of its budget on people and skills that are not the constraint.
- **Do not coach Connor Blake yet.** His financial exposure is real, but causal confidence is not established. Investigate first (territory, deal mix, lead quality) before risking a result that may not be behaviour-driven.
- **Do not change the sales methodology.** The evidence points to inconsistent execution of foundational behaviours, not a flawed approach.
- **Do not use attainment alone to set coaching priority.** Elena's 91% attainment reads as moderate underperformance on its own. Her execution score, the sample's weakest, is what actually tells you where to look first.

From diagnosis to a measured result

0–30 DAYS

Correct Discovery Discipline

- Elena and Marcus into intensive coaching: talk time under 65%, mandatory discovery-depth sequence, weekly manager inspection against that specific criterion.
 - Open the Connor Blake investigation (territory, deal mix, lead quality) before any coaching decision is made about him.
 - Team-wide call standard addressing permission-seeking language, using Priya's calls as the internal reference recording.
-

31–60 DAYS

Demonstrate Behavioural Consistency

- Ten-call re-score for Elena and Marcus under this same methodology.
 - Track execution-score movement and protocol pass rate against the Day-0 baseline, not general impressions from call listening.
 - Close the Connor Blake investigation with a documented conclusion.
-

61–90 DAYS

Prove Commercial Impact

- Compare, before and after, for every coached seller: execution score, attainment, next-step compliance.
- This is a closed-loop measurement, not a one-time correction. Re-auditing against this baseline is what turns a coaching effort into evidence of what worked.

12. LEADERSHIP VERDICT

The one-page answer

COMMERCIAL POSITION

Healthy: 101.6% of target this quarter.

BIGGEST RISK

\$61,400 of quarterly revenue currently sits on the team's least-explained execution profile.

MOST IMPORTANT PERSON

Priya Anand: the only seller proving the standard is achievable here, on this team, this quarter.

EXECUTION POSITION

Below standard: team average 70.3/100 against a 100-point benchmark, spread 55–91.

BIGGEST OPPORTUNITY

Approximately \$241,000 annualised in modelled opportunity, concentrated in two sellers and two behaviours.

ONE MANAGEMENT ACTION

Put Elena and Marcus into structured coaching this week, using Priya's calls as the reference standard. Open the Connor investigation in parallel.

THIS WAS A SAMPLE

See what this looks like for your team.

The Revenue Execution Audit starts with your reps' real calls, not a template. Same maths, same discipline, same standard of proof, built from your data.

Start Your Revenue Execution Audit

www.thexdrcoach.com.au/pricing

One-time diagnostic · Required before Revenue Governance · Typical turnaround 5–7 business days



Meridian Cloud Solutions is a fictional company created solely to illustrate this report's format. Sample Report.