

# The AI Baseline Report

Where AI moves revenue, cost, and risk in this business

## Prepared for the CEO and leadership team

Client anonymized. Figures altered to preserve confidentiality.  
Structure and register unchanged.

---

|               |                                                            |
|---------------|------------------------------------------------------------|
| PREPARED BY   | Attain, Independent Fractional Chief AI Officer Advisory   |
| ARTIFACT TYPE | Report, produced once at diagnostic and refreshed annually |
| DATE          | Sanitized sample, 2026                                     |

**Confidential.** A readiness score tells a CEO how the company compares to a maturity model built for a conference stage. It does not tell them what to decide on Monday. This report contains no score. It states, in the company's own numbers, where AI would move revenue, where it would remove cost, and where it already creates risk, so that the 90-Day Decision Agenda that follows it is an argument about evidence, not opinion.

*This sample contains the cover, the method, the section map of the full report, and one worked finding page. The full report runs 12 to 15 pages and is delivered with the 90-Day Decision Agenda as the output of the two-week diagnostic.*

## THE INSTRUMENT

# Method

---

2 – 6 weeks, 3 data sources, cross-checked.

### 01 The data pull.

12 months of transaction, quote, expense, and system data, requested in writing against a fixed list. For this client: the full quote log, ERP order history, reason codes on lost business, expense lines matched against an AI vendor list, and the item-master. The data was gathered before the first interview so that every conversation starts with a number in context.

---

### 02 The interviews.

11 role-specific interviews, run from written guides: the CEO, the leadership team, and the roles who assemble the quotes, chase the suppliers, and re-key the orders. Leadership describes the process as designed. Operators describe the process as run.

---

### 03 The Shadow AI Audit.

Run inside the same 2 weeks: an anonymous amnesty survey of every employee plus a technical sweep of logs, grants, and expense lines. For this client, 87 of 118 employees completed it. The audit tells where AI operates in the business without permission, which is the signal of baseline demand.

#### WHERE THE SOURCES DISAGREE

The sources never fully agreed. Where the client's data could not support a number, we did not substitute an industry benchmark; the missing number is recorded as a measurement gap with an owner and a cost to fix. Several of this client's most consequential findings are gaps of exactly that kind.

## CONTENTS

# Section map of the full report

---

### 1. The operating picture.

One page: how this business makes and loses money today, in its own figures. The CEO signs this page before we proceed. If leadership does not recognize the company on page one, nothing built on it deserves their trust.

---

### 2. Revenue findings.

Where money enters, at what speed, and what pattern the wins and losses follow. For this client, the central finding was cycle time: won quotes left the building in a median of 4 days, lost quotes in a median of 11.

---

### 3. Cost findings.

Hours counted at the process level, not estimated at the department level. Each finding is a named workflow with an annual hour count and its derivation shown, so the number survives a CFO's scrutiny.

---

### 4. Risk findings.

What the Shadow AI Audit surfaced, single-person dependencies, and the data classes already leaving the building on personal accounts, stated with counts and account types rather than adjectives.

---

### 5. The measurement gaps.

Every number the business could not produce, in one table with owner, cost, and duration. For this client the largest was reason-code coverage: only 38 percent of lost quotes carried any reason at all, which caps what anyone, including us, can claim about win rates.

## 6. AI gaps and opportunities.

The candidate opportunities the findings surface, each paired with the specific data, system, skill, or governance gap that blocks it. For this client the baseline produced 31 candidates; 4 survived scoring. This section carries the summary table; the full scoring, gap register, and declined-opportunity record are delivered in the Gaps & Opportunities Analysis that builds on this report.

---

## 7. Recommendations.

The positions this baseline supports, stated so that someone can disagree with them: what to start this week at zero spend, what to fund first and in what order, and what to decline with the reason on record. For this client: enforce reason codes at point of loss immediately, remediate the item-master before any quoting tool is evaluated, and hold all other AI spend until both are underway. Each recommendation names its owner and feeds the 90-Day Decision Agenda, where the spend arguments are made in full.

---

# What we would not do

---

We would not grade this company against a five-level maturity scale, because a grade produces a feeling and a baseline must produce decisions. We would not fill a missing client number with a vendor case study or an industry benchmark; a report that mixes the client's numbers with borrowed ones teaches the reader to trust neither. And we would not name tools in this report. The baseline states facts. The spend arguments arrive in the decision agenda, and tools enter later, through the evaluation rubric, after the gaps that would sink them are priced.

**A grade produces a feeling. A baseline must produce decisions.**

# Cost Finding 1 of 5

**Quote assembly, standard quotes.** From the baseline of a \$42M specialty distributor, 118 employees. This finding became Decision 2 on the client's 90-Day Decision Agenda and Opportunity 1 in the Gaps & Opportunities Analysis.

|               |                 |                |              |
|---------------|-----------------|----------------|--------------|
| QUOTES A YEAR | STANDARD QUOTES | MIN. PER QUOTE | HOURS A YEAR |
| 5,200         | 3,100           | 79             | 4,100        |

**The count**

The business issues roughly 5,200 quotes a year. 60 percent are under \$25,000 and follow one pattern: an emailed RFQ, a manual match against the item-master, re-keying into the ERP, a priced quote out. Timed across the inside-sales team and reconciled against the quote log, the matching and re-keying averages 79 minutes per standard quote. Across 3,100 standard quotes, that is 4,100 hours a year of senior commercial staff doing transcription.

**The number behind the number**

The hours are the visible cost. The finding that matters is speed: the median won quote went out in 4 days, the median lost quote in 11. That comparison rests on the 38 percent of lost quotes carrying a reason code, so the report states it as a pattern, not a proof, and logs the coverage gap in Section 5 with a zero-spend fix: reason codes enforced at point of loss, 15 minutes of sales-manager attention per week.

|                  |                   |                      |
|------------------|-------------------|----------------------|
| MEDIAN WON QUOTE | MEDIAN LOST QUOTE | REASON-CODE COVERAGE |
| 4 days           | 11 days           | 38%                  |

**Cross-check**

The Shadow AI Audit confirmed the pressure from the other direction: the heaviest unsanctioned AI use in the company sat inside this exact workflow, where inside-sales staff had built their own quote-drafting relief out of personal chatbot accounts. When a process shows up in both the hour count and the shadow traffic, the demand is not hypothetical. The staff have already voted.

## POSITION

This is the strongest cost finding in the report, and it is not primarily a cost finding. Cutting the 4,100 hours pays for tooling; closing the 7-day gap between won and lost quotes is where the revenue case lives, and that case cannot be proven until reason-code coverage rises. The coverage fix costs nothing and starts this week regardless of anything else this diagnostic recommends. Whether and how to fund the automation itself is Decision 2, and it is argued in the decision agenda, not here.

---

*End of sample. The Baseline Report and the 90-Day Decision Agenda are the two deliverables of the standalone diagnostic: two weeks, \$3,000, credited against the first retainer month. The report's findings feed the Gaps & Opportunities Analysis and the 12-Month Roadmap directly.*

**ATTAIN.**

*Attain is an independent AI advisory for founder-led companies.  
Advice is the product. Nothing else is for sale.*