

Cadenta

cadenta.example · Base Cycle · August 2026

**Every question
your crawl asks
is written from this
document.**

Read it as the brief for the crawl, not as a
summary of one. What you correct here is
what we ask the engines.

- 1· Foundations
- 2· Voice of Customer
- 3· What happens next
- 4· Method & sources

CONTENTS

What's in this review

The map your crawl runs on: who your buyers are, what we'll ask the AI engines on their behalf, and how both were built. Every question we send out comes from these rows, which is why you see them first. Sections are numbered on every heading, so 1.4 is where the contents says it is.

1	Foundations	4
1.1	Who buys this	5
1.2	Who you're compared against	5
1.3	What buyers evaluate	6
1.4	Why buyers search	7
1.5	AI crawler access	8
1.6	Content quality signals	8
1.7	Findings to fix	9
2	Voice of Customer	12
3	What happens next	14
3.1	How your questions get written	14
3.2	What the crawl adds	15
3.3	Make it sharper with your sales calls	15
4	Method & sources	17
4.1	How your foundation was built	17
4.2	How the crawl will measure you	18

SUMMARY

What we found

3 of your 6 findings are high severity. Those are the ones that change whether an AI engine can read and cite you at all; the rest are worth doing after them.

FINDINGS	WEAKEST SIGNAL	CRAWLER ACCESS
6 3 high · 2 medium · 1 low	0.39 Content freshness	4 of 4 AI answer-engine crawlers can reach you

FIX THESE FIRST

01 Pricing page renders its tiers client-side, so crawlers see an empty table

HIGH

ENGINEERING

1-3 DAYS

02 Documentation sits behind a login wall

HIGH

ENGINEERING

1 WEEK

03 The product page is published at two URLs with no canonical

HIGH

ENGINEERING

< 1 DAY

3 further findings — 2 medium · 1 low — with the same detail on each. [all 6 findings → · page 9](#)

Content freshness. Most scored pages haven't changed in months. Engines favour recent sources, so this is the cheapest lift you have. [all 4 signals → · page 8](#)

What isn't here yet. Nothing in this document measures how often an AI engine actually recommends you — that is the crawl, and it has not run for this cycle. What is here is the map it runs on and the technical state of your site. [what the crawl adds → · page 15](#)

SECTION 1

1

Foundations

An engine recommends whoever best fits the buyer it thinks is asking. So before we measure anything we write down who that buyer is, who you get compared to, and what they weigh — plus what a crawler can and can't read on cadenta.example today. Every question in this crawl comes from these rows, which is why you approve them first.

BUYERS 5 personas	COMPETITORS 6 in your category	FEATURES 8 buyers evaluate
PAGES 48 read of 93 found	FINDINGS 6 3 high · 2 medium · 1 low	

Why this section comes first. An AI engine can't recommend you until it can describe you, and it describes you using what it can read: your site, your competitors' pages, and the sources your category already trusts. So that is where we start, on purpose. Modelling your market from public evidence is fast and repeatable, and it puts a defensible map in front of you in days rather than the weeks a research project would burn.

Then we layer. What the engines already say about your category, and how your competitors actually position, are read into these rows before you see them. Your sales calls are the layer after that — the one that moves the ordering from what the category suggests your buyers weigh to what they were recorded saying, and the reason nothing below is attributed to a named buyer yet. The next section shows exactly what changes when they go in.

None of this is neutral description. Choosing which buyers count, which rivals you are measured against and which capabilities matter is choosing how you enter the arena — and every question in this crawl was written from these rows. Read it as the positioning call it is, and tell us where it misses.

1.1

Who buys this

5

The people who have to be convinced, and the ones your buyer questions get written for.

BUYERS 5 in the buying group we model	CAN BLOCK THE DEAL 3 of 5 hold a veto — convincing the others is not enough
--	--

Engineering managerMANAGER

Engineering

Operations leadDIRECTORCAN BLOCK THE DEAL

Operations

Agency ownerFOUNDER

Leadership

IT and security reviewerSENIORCAN BLOCK THE DEAL

IT

Finance approverCONTROLLERCAN BLOCK THE DEAL

Finance

1.2

Who you're compared against

6

The vendors a buyer weighs you against, most directly compared first.

01

Kanbrix

CATEGORY LEADER

The default board tool, sold on how fast a team can start using it.

02

Sprintwell

DIRECT

Engineering-first planning with the deepest issue tracker in the category.

03 Northpath ENTERPRISE

Consulting-led rollouts for regulated organisations that buy on governance.

04 Werkstream AUTOMATION

Workflow automation first, project views second.

05 Momentro LIGHTWEIGHT

Minimal, opinionated, priced per workspace rather than per seat.

06 Pathgrid ADJACENT

Resource and capacity planning expanding into delivery tracking.

1.3 What buyers evaluate ⁸

The capabilities buyers weigh when they compare you, ordered by how much of your question set each one earns.

CAPABILITIES

8

2 strong · 3 moderate · 3 weak

WHERE YOU'RE EXPOSED

3 of 8

rated weak or absent — highest-ranked is #1, Integration coverage

01 Integration coverage WEAK

Does it actually talk to the tools we already run?

02 Reporting and rollups WEAK

Can I get a portfolio view without exporting to a spreadsheet?

03 Migration path WEAK

How much work is it to move four years of history across?

04 Permissions model MODERATE

Can a client see their project and nothing else?

-
- 05 **Security and compliance** MODERATE
Are you SOC 2, and will your DPA survive our legal review?
-
- 06 **Pricing transparency** STRONG
What does this cost at forty people, without a sales call?
-
- 07 **Automation rules** MODERATE
Can it move the ticket itself when the build passes?
-
- 08 **Onboarding effort** STRONG
Will the team still be using it in three months?
-

1.4 Why buyers search ⁵

The problems that put a buyer in the market, ordered by severity.

PAIN POINTS

5

2 high · 2 medium · 1 low

-
- 01** **Work is spread across four tools and no one view is true** HIGH
Status lives in three places and none of them agree.
-
- 02** **The last rollout was abandoned within a quarter** HIGH
We bought one of these already. Nobody opened it after March.
-
- 03** **Reporting to leadership is still assembled by hand** MEDIUM
I rebuild the same deck every Monday morning.
-
- 04 **Per-seat pricing punishes bringing in contractors and clients** MEDIUM
I'm not paying full price for someone who logs in twice a month.
-
- 05 **Nobody can estimate the migration cost with a straight face** LOW
What happens to four years of tickets?
-

1.5 AI crawler access 4 / 4

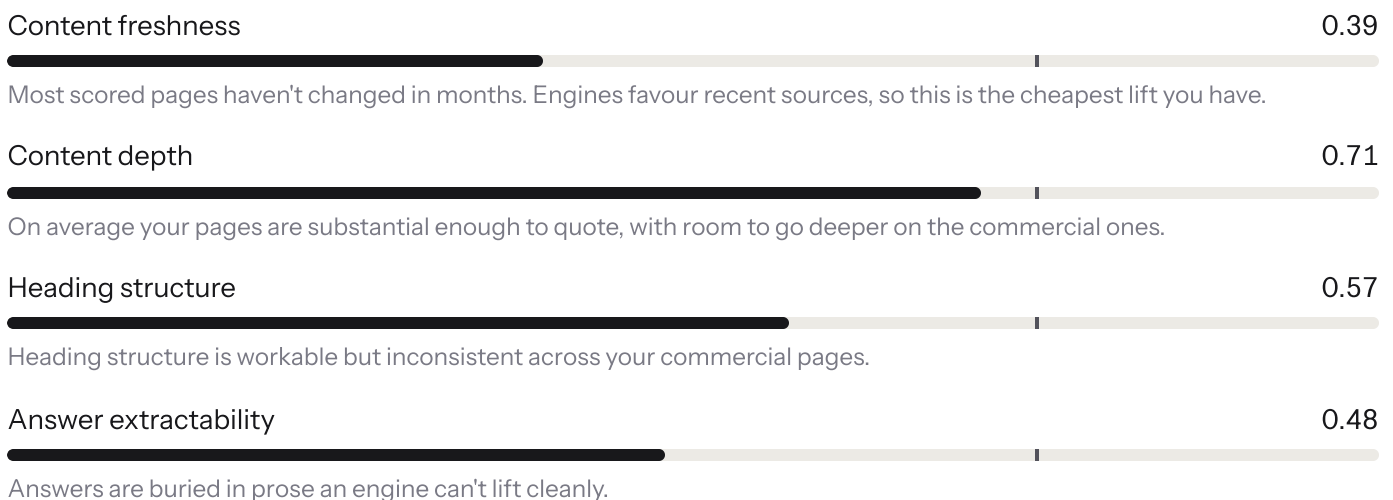
Every crawler that feeds an AI answer engine can reach cadenta.example. CCBot is blocked, but it does not feed one — GPTBot, ClaudeBot, PerplexityBot, Google-Extended are each allowed by your robots.txt.

CRAWLER	STATUS
CCBot	blocked
GPTBot	allowed
ClaudeBot	allowed
PerplexityBot	allowed
Google-Extended	allowed

1.6 Content quality signals

AI platforms preferentially cite fresh, well-structured, deeply-covered pages they can extract clean passages from. Here's how the pages we read score.

Each bar is one score averaged across the 48 pages we analysed. The line at 0.75 is where we treat that average as healthy; below 0.50 we treat it as holding citations back. Individual pages can sit well below their own average — where a specific page or group is the problem, it is written up in the findings with its own number.



1.7 Findings to fix ⁶

Ordered by severity — high first. Each one is a concrete fix, with an owner and an effort estimate.

#	FINDING	SEVERITY	OWNER	EFFORT
01	Pricing page renders its tiers client-side, so crawlers see an empty table	HIGH	Engineering	1–3 days
02	Documentation sits behind a login wall	HIGH	Engineering	1 week
03	The product page is published at two URLs with no canonical	HIGH	Engineering	<1 day
04	Placeholder copy is live on two commercial pages	MEDIUM	Content	1–3 days
05	No integration page states what the integration actually does	MEDIUM	Content	1 week
06	Two changelog entries share a title tag	LOW	Content	<1 day

01 • Pricing page renders its tiers client-side, so crawlers see an empty table

HIGH

ENGINEERING

1–3 DAYS

Rendering

Owner · Engineering

Effort · 1–3 days

WHAT WE FOUND

The plan comparison is hydrated from a JSON call after load. In the raw HTML the table has headers and no rows.

WHY IT MATTERS

Price is the most-asked question in this category and the one buyers ask an assistant precisely because they do not want a sales call. Served as an empty table, the answer goes to whoever publishes a number.

RECOMMENDED FIX

Server-render the tier table and let the interactive toggle enhance markup that is already in the HTML.

Scope: The pricing page and the two comparison pages that embed the same component

02 • Documentation sits behind a login wall

HIGH

ENGINEERING

1 WEEK

Access

Owner · Engineering

Effort · 1 week

WHAT WE FOUND

All 340 documentation pages redirect anonymous requests to the sign-in screen, including the integration directory.

WHY IT MATTERS

Docs are what an engine quotes when a buyer asks whether a tool supports something specific. Behind a login they are invisible, and competitors' equivalents are not.

RECOMMENDED FIX

Open the reference and integration docs to anonymous readers and keep only account-specific pages gated.

03

● The product page is published at two URLs with no canonical

HIGH

ENGINEERING

< 1 DAY

Duplication

Owner · Engineering

Effort · < 1 day

WHAT WE FOUND

/product and /platform both return 200 with identical content and neither declares a canonical.

WHY IT MATTERS

Every citation and internal link splits across two addresses, so neither accumulates the authority one would.

RECOMMENDED FIX

301 /platform to /product and add a self-referencing canonical.

04

○ Placeholder copy is live on two commercial pages

MEDIUM

CONTENT

1-3 DAYS

Content quality

Owner · Content

Effort · 1-3 days

WHAT WE FOUND

A lorem-ipsum block and a heading reading “TEMPLATE — replace before launch” are live on the manufacturing and nonprofit landing pages.

WHY IT MATTERS

Placeholder text is a strong low-quality signal, and it sits on pages built to be cited.

RECOMMENDED FIX

Replace or unpublish both pages, then add a pre-publish check for placeholder strings.

05

○ No integration page states what the integration actually does

MEDIUM

CONTENT

1 WEEK

Depth

Owner · Content

Effort · 1 week

WHAT WE FOUND

The 62 integration pages carry a logo, a one-line blurb and a Connect button. None describe which objects sync, in which direction, or how often.

WHY IT MATTERS

“Does it sync two-way with our issue tracker” is a question engines answer from documentation. With nothing to lift, they answer it about somebody else.

RECOMMENDED FIX

Give each integration page a sync-direction table, an object list and a stated refresh interval.

06

○ Two changelog entries share a title tag

LOW

CONTENT

< 1 DAY

Metadata

Owner · Content

Effort · < 1 day

WHAT WE FOUND

Two 2025 posts both title as “Product update”.

WHY IT MATTERS

Duplicate titles make it ambiguous which page answers which question.

RECOMMENDED FIX

Retitle each to the change it actually describes.

SECTION 2

2

Voice of Customer

What your buyers say in their own words, validated against every row of section 1 — built from your recorded sales calls. This cycle ran on category research alone, so this page is a preview of what the recordings add.

When you share recorded calls, here is what this section holds:

What we heard

The patterns across your calls — each claim carrying the quotes that back it.

Who said it

The functions actually in the room, counted by accounts and people, set against the buyers your foundation models.

The pushback, in their words

Objections grouped by what was said, rather than by what we modelled.

What they ask you for

The requests that repeat across accounts.

What we'll optimise for

What the conversations change about the questions we run — row by row, each one linked back into your foundation.

SPECIMEN • AN INVENTED FINDING, DRAWN THE WAY YOUR REAL ONES ARE

Buyers hear “integration” and ask whether the sync is two-way

4 ACCOUNTS

The site says integrations; the calls say buyers only trust the ones that write back.

“If it can't push changes back, it's an import, not an integration.”

A buyer on one of your calls • 00:31:08 • linked to the recording

Only your buyers' and customers' own words count — never ours, never a partner's — and every quote links to the moment in the recording, so any claim can be checked against the conversation rather than taken on trust. Share the recordings and the next cycle builds this section from them, revalidates section 1 against it,

and re-weights the questions we run before anything is measured.

SECTION 3

3

What happens next

What this document turns into, what it costs to get a row wrong, and how to make it sharper before the crawl starts.

3.1 How your questions get written

This is the input, not the summary. A visibility crawl is only as well-aimed as the map it is written from, and this document is that map.

YOU ARE HERE

This document

5 buyers, 8 capabilities and 5 pain points, each one a named row you can correct.

NEXT

Your buyer questions

We write the prompts a real buyer would put to an AI engine — one line of questioning per buyer, aimed at the capabilities they weigh and the problems that put them in the market. Every question comes from a row above, and carries which one.

THEN

Your visibility

We put that set to ChatGPT, Claude, Gemini, Perplexity and record what comes back — whether you appear, whether you are recommended, and who is recommended instead. Because each answer traces to a question and each question to a row, the result reads back onto this document rather than beside it.

What a wrong row costs. It does not produce a wrong answer later — it produces a question nobody asks. We would measure your visibility accurately, against prompts your buyers never type, and every number after that would be true and useless. That is the one failure this review exists to prevent.

Three things worth checking as you read Foundations:

- **A buyer who isn't in Cadenta's market.** Every persona becomes its own line of questioning, so one that doesn't buy from you spends a share of the crawl on somebody else's category.
- **A capability we've ranked too low.** The order in §1.3 is the order we write to. Something you win on, sitting near the bottom, is a question set that never asks about your best argument.

- **Language that isn't how your buyers talk.** The plain-English gloss under each row is what the questions are phrased from. If it reads like your marketing rather than like their search, tell us — that gap is most of the difference between appearing and not.

3.2 What the crawl adds

Once this foundation is signed off, the report grows three sections around it. Your copy is reissued each time one lands, so the file you are holding is always the whole of what exists.

Questions	Every prompt we put to the engines, grouped by the buyer asking it, with the capability and pain point each one probes. Written from the rows in Foundations.
Visibility	Where you stand across that set — how often you appear, how often you are recommended, your share of voice against the competitors in §1.2, and which sources the engines cite instead of you.
Action plan	The ranked work, in execution order, each item naming the questions it would win and the pages it touches.

3.3 Make it sharper with your sales calls

Foundations is built to be layered, and it is one layer short. We start from public evidence because it is fast and repeatable: your market, your site, your competitors' positioning, checked against what the engines already say. That is enough to write a question set that finds where you're losing, and it is why you had a map in days.

The layer it doesn't have is your buyers in their own voice. Public evidence can tell us what your category weighs; only your calls can tell us what *your* accounts actually raised, in what order, in what words. If you record them, we read them, and this document changes in three specific ways.

TODAY

- Rows are ordered by modelled priority and severity.
- Nothing is attributed — no row can name who raised it.
- Buyer language is written from your market, not quoted from it.

WITH YOUR CALLS

- Rows are ordered by how many of your accounts actually raised each one.
- Each row carries the count, and the top rows carry a verbatim quote with the speaker and the moment in the recording.
- Objections and buying triggers come out of the calls rather than the category — including the ones we would never have modelled.

The window closes at sign-off. The question set is pinned to the foundation you approve, so calls shared after that go into the *next* cycle rather than this one. This foundation is already signed off — anything you share now sharpens the cycle after it.

Two ways in, whichever is less work:

- **Upload what you have.** Call exports, transcripts or notes, at resonatelabs.co/account/transcripts/. Formats are broad; if it is text of a real conversation, it is usable.
- **Hand over your notetaker.** Read access to Gong, Fathom, Grain or Otter and we pull the corpus ourselves. Nothing is written back, and nothing is shared outside your account.

What comes out of it lands at resonatelabs.co/account/voice-of-customer/ — what buyers came looking for, what they asked you for, and what stopped them, each ranked by how many accounts said it.

SECTION 4

4

Method & sources

Where every line of this came from, so it can be checked before you sign it off.

4.1 How your foundation was built

We build this from category research: your own site, the sites of the vendors buyers weigh you against, how the category is listed and described publicly, and what reviewers say on the public review platforms. You then review it and correct anything we got wrong, and nothing is measured until you approve it.

That makes it our reading of your market, checked by you — not a claim about what your buyers said, because on this cycle we have not heard from them directly. Recorded sales calls are how that changes. Here is what it looks like, on an invented row: the same line of section 1, before and after your buyers' own words go in.

THIS CYCLE · CATEGORY RESEARCH, REVIEWED BY YOU

Renewal arrives before anyone can prove it worked

HIGH

RESEARCH ONLY

The champion who bought it is left defending the spend with anecdotes.

NEXT CYCLE · THE SAME ROW WITH YOUR CALLS IN

Renewal arrives before anyone can prove it worked

HIGH

3 ACCOUNTS SAID THIS

PRIORITY #1

The champion who bought it is left defending the spend with anecdotes.

"We had no number to put in front of finance at renewal, so it got cut."

A buyer on one of your calls · 00:14:22 · linked to the recording

What your buyers confirm stays. What they contradict is changed. What they raise that the research missed is added — every quote carrying the speaker, the account and the second it was said. Share the recordings and section 1 is rebuilt against them on the next cycle, before anything is measured.

Sources we read directly:

- Your site — [cadenta.example](#) (48 of 93 pages read)
- Your sitemaps — the sitemap index and all 4 child sitemaps (131 URLs, 38 URLs excluded)
- A live crawler reachability check — 5 crawler user-agents requested your homepage directly, rather than reading your robots.txt alone

4.2 How the crawl will measure you

Once the questions are written, each one is put to every AI engine we run and the full answer is recorded along with the sources it cited. Visibility is the share of questions where an engine puts you in the answer. Win rate is graded, not binary: a solo top pick scores 100%, a strong second 50%, a place in a list 25%, a passing mention 10%, and absence 0 — averaged across every question, including the ones you are absent from. None of that has run for this cycle; it is stated here because it is what your foundation is about to be turned into.

The technical findings come from a crawl of 48 of the 93 pages we could discover on your own site. Archive, tag and legal pages are excluded from that count, and because the site holds more pages than one crawl reads, the crawl always includes your most recently updated pages alongside every commercially relevant one, up to the crawl's page budget. Each page is checked for crawler access, freshness, structure and how cleanly a passage can be extracted and quoted.

Crawl	cadenta-base
Data as of	July 24, 2026
Questions	150
Engines	ChatGPT, Claude, Gemini, Perplexity
Answers captured	600



This foundation review was produced by Resonate Labs.

We measure how AI engines answer the questions buyers actually ask — across ChatGPT, Claude, Perplexity and Gemini — and hand back what to change, in the order it is worth changing.

Your foundation is ready for your review. The crawl is what turns it into numbers — where you actually stand across the full question set and every engine we run. Reply to the email this came with and we'll start it.

resonatelabs.co

