



PRODUCT UPDATE · SEPTEMBER 2026

AXAT Assessments **are live.**

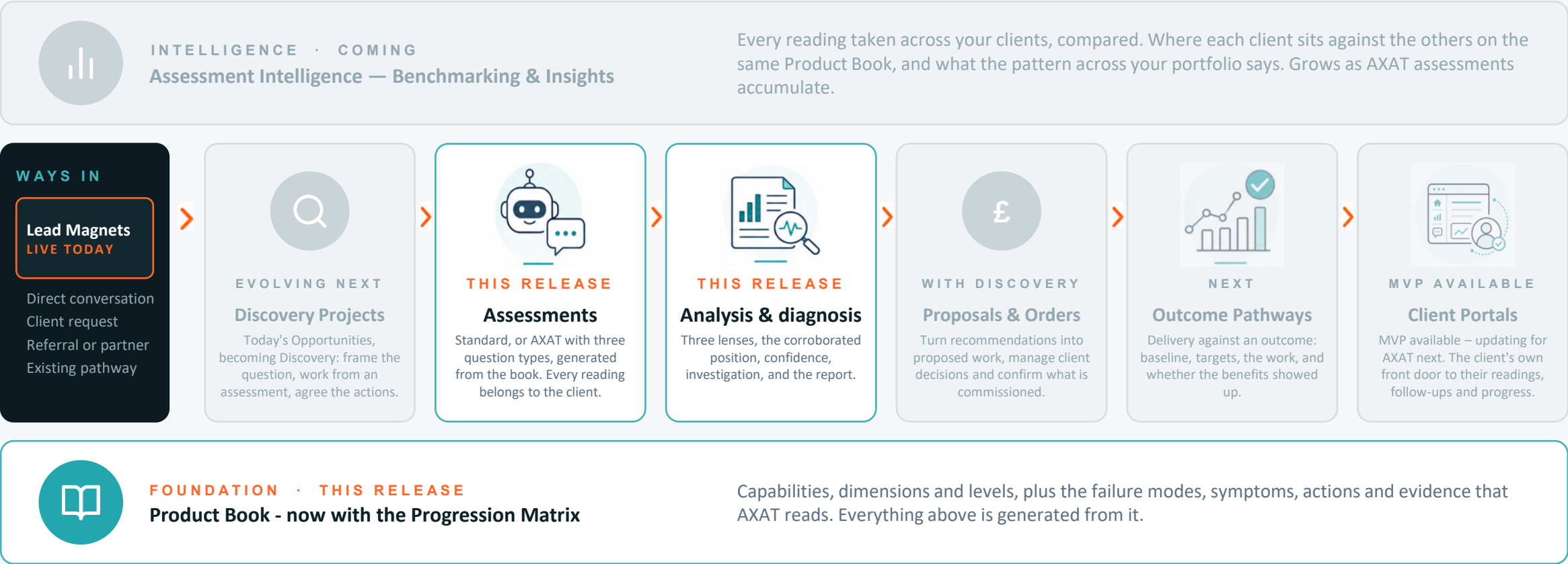
A quick guide to what's new in TheAX - and what to do first to use the new AX Assessment Triangulation (AXAT) Assessments.

Turn “trust us, it worked” into a number - the first triangulated diagnostic method, now in the assessments you already run.

Everything you already have in TheAX, your templates and your Product Books all continue exactly as before.

TheAX has a new assessment engine: **AXAT**.

TheAX platform has been updated with an enhanced Assessment engine - AXAT (AX Assessment Triangulation). This updates the three areas highlighted below. How these have been updated are briefly explained in this deck.



Assessments ask for a score. **AXAT asks whether it is true.**

AXAT is an assessment method and algorithm developed by TheAX. It elevates maturity assessment from a useful tool to a repeatable and defensible audit system that delivers your best diagnostic thinking repeatedly no matter who runs it. Every capability is read through three independent lenses. Opinion stays useful as input - it just stops being mistaken for the conclusion.

A TYPICAL ASSESSMENT

"Where would you rate yourselves?"

3.8

self-reported maturity

A great way to get fast assessment for client engagement and conversations.

Here the claim becomes the result. Optimism, politics, recency and interpretation are averaged into one number. That reduces the repeatability over time.



AXAT · THREE INDEPENDENT READINGS

Perception

what people believe

3.8

Evidence

what can be shown

2.4

Symptoms

what people experience

2.6

Corroborated position **2.4** 84% confidence · open questions retained

- ✓ The lowest lens sets the position – and the gaps show you where to focus improvement.
- ✓ Confidence is reported beside it, together with what would make the reading stronger.
- ✓ Nothing is averaged away – readings are repeatable and defensible.

Two restaurants, one street. The difference is underneath.

A rating tells you there is a gap. The pattern across the three lenses tells you what kind - and therefore what to do.

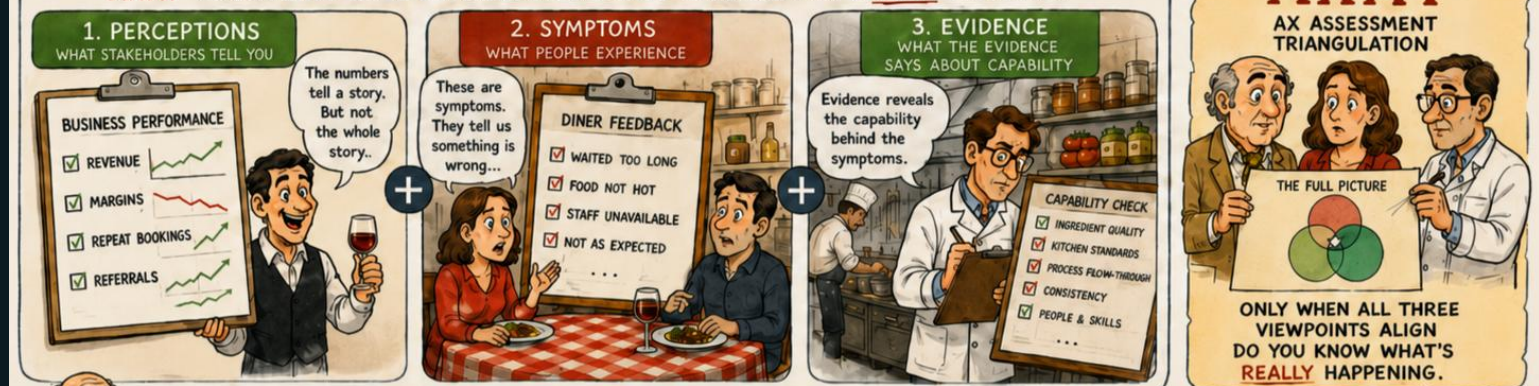
Same kitchen, same suppliers, very different results. Mario's diners are happy and nothing is written down. Gina's kitchen is built and the service isn't landing. One set of questions can't tell them apart.

How AXAT delivers faster, higher-quality assessment diagnostics

WHY DO TWO ITALIAN RESTAURANTS GET VERY DIFFERENT RESULTS?
THE MENU LOOKS FINE. THE KITCHEN LOOKS FINE. THE DIFFERENCE IS UNDERNEATH.



MOST ASSESSMENTS ONLY ASK ONE SET OF QUESTIONS AND HOPE FOR THE TRUTH.
AXAT TRIANGULATES THREE DIFFERENT PERSPECTIVES TO REVEAL WHAT'S REALLY HAPPENING.



AXAT CONNECTS PERCEPTIONS, SYMPTOMS AND EVIDENCE TO UNCOVER THE CAPABILITY THAT DRIVES PERFORMANCE AND SUSTAINABLE IMPROVEMENT.

MOVE BEYOND OPINIONS.
BUILD EVIDENCE. DRIVE REAL IMPROVEMENT.

The **AX**.ai
TURNING EXPERTISE INTO SCALABLE SYSTEMS

The same two restaurants, read by AXAT.

Three lenses each. Two failures that look identical on a score and need opposite fixes.

MARIO · THE HERO CHEF

Perception

3.1

Evidence

1.67

Symptoms

3.44

Corroborated at 1.67 - the evidence floor.

The diners are happy, but nothing is written down and nothing would survive Mario taking a week off. Happy diners cannot rescue a kitchen that does not exist.

THE ADVICE Build it and get it out of one head. Write the recipes down, put the prep system in place, train a second pair of hands.

Same score band on a standard Perception based assessment. But opposite advice is needed.

GINA · THE BUILT KITCHEN

Perception

3.5

Evidence

4.0

Symptoms

2.33

Corroborated at 2.33 - the symptoms cap it.

The recipe book, the certificates, the trained staff are all real. The service is not. Evidence alone would have waved her through.

THE ADVICE Stop writing. Make it land. More documentation would have raised the score and changed nothing.

That is the difference between a rating and a diagnosis - and it is the reason AXAT never averages the three lenses together.

A rating tells you there is a gap. The pattern across the three lenses tells you what kind - and therefore what to do.

Three modes now. **AXAT is the one that tests the claim.**

A third assessment mode, alongside Lead Magnets and Standard. Chosen per assessment, once a Product Book is ready for it.

LEAD MAGNET · UNCHANGED

Self-completion

A short assessment for lead capture or quick data gathering.

One view, one person

Starts the conversation with low effort.

STANDARD · UNCHANGED

One question per capability

Where do you think you are? Scores are averaged, so a strength can offset a weakness.

A fixed report

Generated once the data is in.

AXAT · NEW

Three lenses per capability

Perception, symptoms and evidence, from the Progression Matrix.

The lowest lens sets the score

Confidence is reported separately, per capability.

A diagnosis, with provenance

Which mechanism is holding a capability back, and what kind of gap it is.

WHAT THIS MEANS FOR YOU

Findings that survive challenge

A corroborated position you can defend in the room, with the evidence trail underneath it.

A diagnosis, not a number

What kind of gap each capability has, and therefore what kind of work, before a day of delivery is sold.

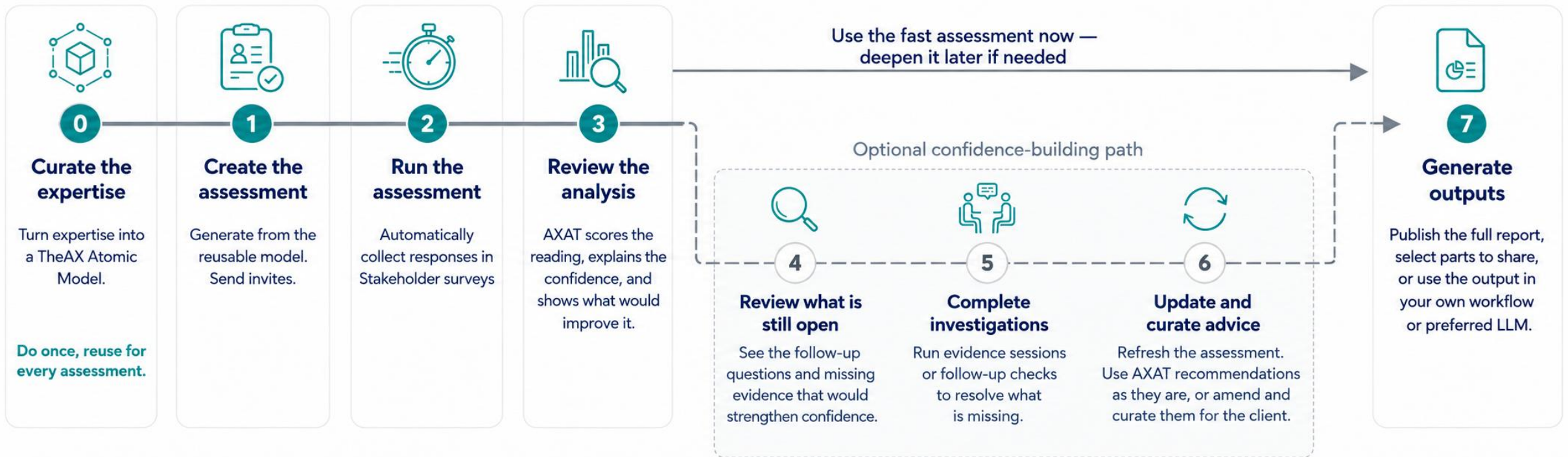
Senior time where it counts

First-pass analysis is done for you. Consultants spend their time on investigation, judgement and advice.

Your Lead Magnets, Standard assessments, templates and Product Books are untouched. AXAT is added, not swapped in.

Run it fast. Then **improve the confidence** where it matters.

Steps 0–3 give you a defended reading. Steps 4–6 are optional for the consultant to lead, and the platform tells them which open questions are worth settling.



The Progression Matrix: what goes wrong, what fixes it.

A new 'Progression Matrix' can be generated using AI for your maturity model to enable AXAT Assessments. This adds depth to the maturity definitions.

So, you get not only what each level of maturity looks like but also, the capability failure modes, how these show up, and how to improve at each level. This is your expertise codified and reusable in ways a playbook never could be.

1. Created by Category, sectioned by dimension

Create only at Category level or with Dimensions that each get their own rows. Read across for progression, down a column for what a level asks for across the failure modes.

2. Rows are failure modes

Each row is one thing that goes wrong (root causes), defined in the context of each level of maturity.

3. Each cell is atomic

One failure mode at one level: one symptom, the actions that address it, the evidence that proves it's done. An atomic map of connections within your area of expertise.

CATEGORY WORKBENCH > Sales Execution

Overview Dimensions Assets Maturity Levels Progression Matrix

The maturity description is the primary driver for all AI-generated content here. Edit it to steer the AI, then save to refresh the affected cells.

Show All Symptoms How to address it Evidence

Sales Activity Management

This dimension measures how well an organisation plans, prioritises, schedules, and tracks sales activities to ensure consistent and effective progression of opportunities through the sales pipeline, prioritisation based on opportunity value and stage, and disciplined follow-up processes to maximise conversion rates and minimise pipeline leakage.

contributes 16 questions Nothing outstanding A way this fails to hold Add failure mode Regenerate

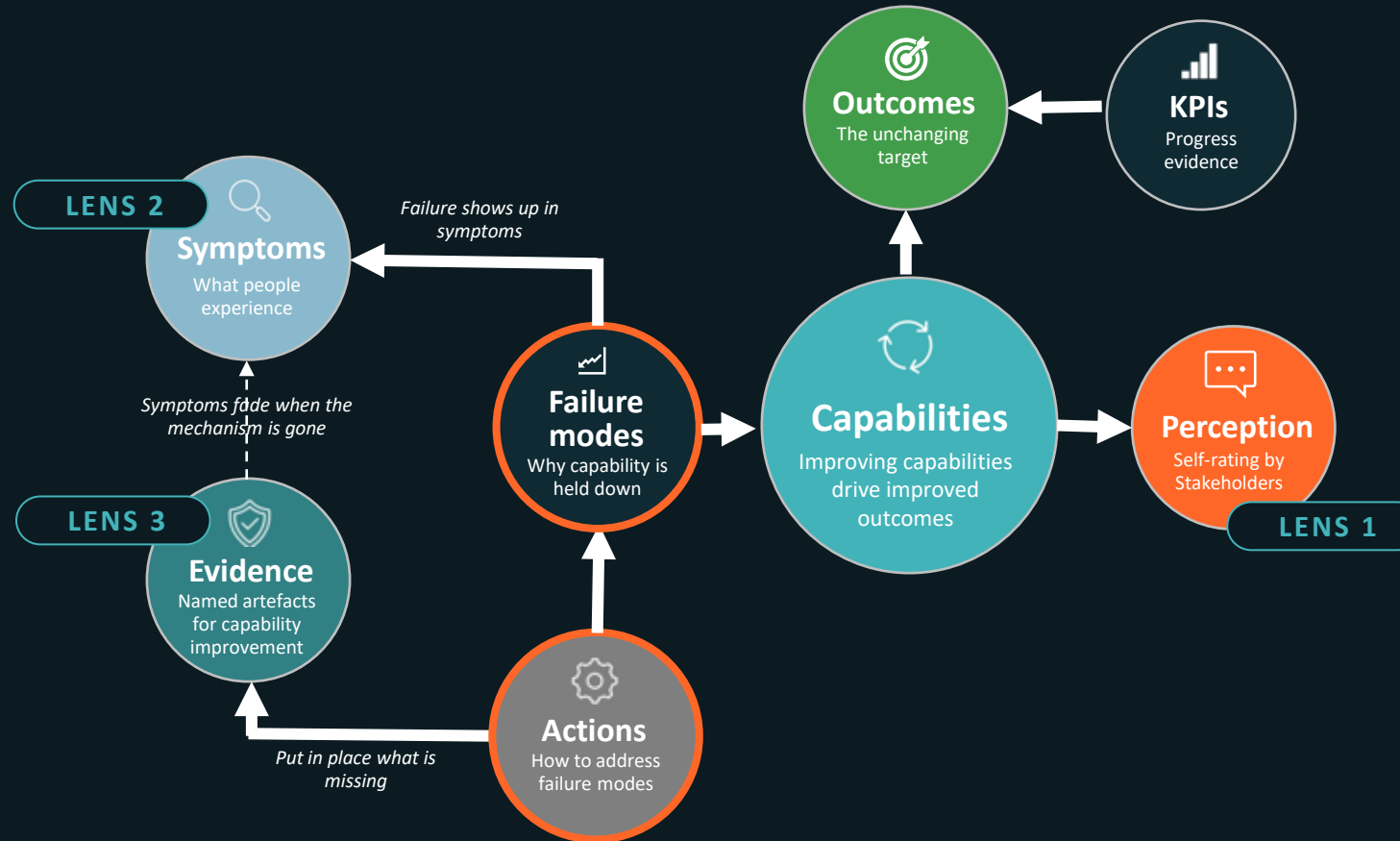
	L1 - Initial Pipeline Setup not asked Sales activities are planned and tracked inconsistently, with limited prioritisation and sporadic follow-up.	L2 - Defined Pipeline Processes 4 questions Sales activities are planned and prioritised with emerging structure, focusing on aligning efforts with pipeline stages and opportunity value.	L3 - Consistent Pipeline Delivery 4 questions Sales activities are planned and executed consistently with prioritisation based on opportunity value and stage, ensuring steady progression through the pipeline.	L4 - Managed Pipeline 4 questions Sales activities are planned and tracked using tools to ensure steady progression and minimise pipeline leakage.
What goes wrong				
Sales activities are not prioritised or tracked, leading to inconsistent follow-up and missed opportunities. L1, L2, L3, L4, L5	Salespeople choose which opportunities to follow up on without a clear plan or tracking, and follow-up is often forgotten or delayed. Introduce a basic template for logging and scheduling sales activities.	Sales teams use a standardised process for tracking activities, but prioritisation is inconsistent and some high-value opportunities are neglected. Implement a prioritisation framework based on opportunity value and stage. Prioritisation criteria document	Sales activities are tracked and prioritised, but some follow-ups are still missed due to lack of automated reminders or oversight. Adopt a CRM system with automated activity reminders and oversight dashboards. CRM activity log with reminder audit	Follow-ups are rare and there are occasional lapses in responding to changes in opportunity status. Integrate real-time monitoring and alerting for changes. Pipeline monitoring
Sales activity plans are created but not consistently executed or	Sales activity planning is ad hoc and rarely documented; plans are not	Sales activity plans are documented but execution is inconsistent, and	Execution of sales activity plans is mostly consistent, but there is limited	Sales activity plan effectiveness are

Workbench > Progression Matrix tab

A capability with dimensions: rows per failure mode, one cell per level

The model connects the lenses to the advice.

The Atomic Model already had the three lenses structured in it. AXAT scores drive the diagnosis and actions required to address capability improvement from identified failure modes.



Perception says where stakeholders think they are. Evidence and symptoms say whether that's true.
The model defines improvement action.

Making a book AXAT-ready.

A maturity model in your Product Book without a Progression Matrix is a perception model. It still runs your Standard Assessments perfectly well. To make your model AXAT ready here are things to consider.

What must exist	What the Workbench checks
Failure mode mechanisms	At least one failure mode per capability (or per dimension)
Symptoms	At least 3 per scored level from L2 up - 4 preferred
Evidence	At least 1 artefact per scored level - a warning below 4
Every mechanism is answered	At least one action or artefact addresses it
Traces are contiguous	No hole in the middle of a failure mode's ladder through the levels.
No near-duplicates	No two items in one lens describing the same observable

AI will generate a Progression Matrix for you based on the rules and what it knows about your expertise area. This is for you and your consultants to check and amend as needed.

Readiness is computed for the whole book, not ticked per item. L1 is the assumed base and carries only information on symptoms and actions for improvement. Aim for at least five failure modes per category.

Workbench › generate the matrix

Per-dimension build flow or Generate all, then validate cell by cell

Build the Progression Matrix

You have written what this capability is. The next step turns that into the questions an assessment asks: what goes wrong, how people notice, what to do about it, and what proves it was done. All of it comes from your definitions, so write those first and this step is quick.

Keep it simple

Build one set of questions for the capability as a whole. A report will tell the client what is holding them back.

Your dimensions stay as description, and are not used here

Go into more detail

Build a set for each dimension. A report will tell the client what is holding them back and which part of the capability it sits in.

More to review now, and more precise findings later

Build the matrix

Takes about 30 seconds.

Either way the client gets one score per capability — this choice only changes how precisely a finding can be pinned down. You can change your mind later, but it means starting this matrix again, because questions written for the capability as a whole do not fit a single dimension.

Book readiness status

AXAT-ready indicator with the rules still needing attention

Product Books / AXcend / Maturity Model

Overview

Files

Proposition

Maturity Model

Outcomes & KPIs

Maps

Assessments

Document Templates

AXcend Maturity Model

Active

Ready to assess

Build the other 1

Maturity Levels →

Categories ↓

Genesis

Consulting is delivered through personal expertise and bespoke problem solving. Value creation is individual, family, and...

Custom-Built

Consulting uses repeatable methods and frameworks, but remains project-centric and...

Productisation

Consulting IP is packaged into defined offers, diagnostics, and programmes. Outcomes are...

Setting up an assessment in **AXAT mode**.

The same wizard you use today, with the mode chosen up front. Per assessment, not per book.

✧ Question Wizard

🔗 Assessment Method

How this assessment establishes its result.

The two methods ask different questions, of different people, and support different claims. This choice decides the rest of the setup.

STANDARD MATURITY

Perception

Captures what people believe about where the organisation stands.

Respondent time2–4 minutes

Panel neededAny size, including one

Consultant inputNone

ResultMaturity profile by category

Reports what is claimed, not what is corroborated. Suitable wherever an unevidenced view is enough.

THEAX-AT

Triangulation

Tests what people believe against what they experience and what actually exists.

Respondent time30–90 minutes, by scope

Panel neededThree or more per capability

Consultant inputAn evidence validation session

ResultCorroborated level, binding lens, confidence

Sets template type, purpose and scope automatically — triangulation needs a panel and produces an authoritative result, so those are not offered as choices. How long it takes depends on how much of the book is in scope, shown once you choose.

▶ Compare the two side by side

⚠ The method cannot be changed once you continue. The two produce different question sets and different setup, so switching later means starting again.

Assessment wizard › mode step

Standard / AXAT selection with the book's readiness shown

1. Choose AXAT when the book is ready

Start as normal. AXAT is only offered on a book that passes the readiness rules, and the method cannot be changed once you continue. Plan for three or more respondents per capability.

2. Section wording stays yours

AXAT Assessments have questions fully drafted for you and organised into sections by category. Each section has wording you can tailor. Respondents can skip sections or questions they can't answer.

3. Watch the question count

AXAT takes a respondent 15 minutes to an hour depending on how many mechanisms are in your Progression matrix, and whether they can answer all sections and questions.

What respondents now see: **three lenses** on one claim.

Three kinds of question per capability instead of one.



Perception

the claim · unchanged

“Where would you place this capability today?”

One answer per category. Routes the rest of the questions for levels immediately at or around their answer level.



Symptoms

the felt pain · new

“Confusion or overlap in sales roles causes delays.” Never ... Always

Behaviour, not documents. A symptom that persists means the level hasn't consolidated, whatever the paperwork says.



Evidence

the proof · new

“Does a documented set of sales roles exist?” Fully in place · Partial · We do not have this · Don't know

One artefact per question. A single respondent confirming it is enough - what matters is that it exists.

New AXAT Test 17-09

PIPELINE FOUNDATIONS

Which of these best describes how this works here today?

Pick the one that comes closest. None of them will describe it exactly.

- Basic sales roles and processes are emerging with inconsistent CRM use, leading to fragmented pipeline management.
- Sales roles and processes are defined and documented, with CRM use becoming consistent to support pipeline stability.
- Sales roles, processes, and CRM use are clearly defined and consistently applied to ensure a stable and repeatable pipeline foundation.
- Pipeline foundations are robust with clearly defined roles, standardised processes, and consistent CRM use, supported by data and governance to ensure stable pipeline management.
- Pipeline foundations are fully embedded, with adaptive roles, processes, and CRM use ensuring a scalable, reliable, and continuously optimised sales pipeline.

Previous Next

PIPELINE FOUNDATIONS

How often does each of these happen?

Think about the last three months — rarely means once or twice, sometimes means monthly, often means weekly. If something does not apply to your part of the business, choose Not sure rather than guessing. You can add detail to any answer; comments may be quoted in the findings, so leave out anything you would not want repeated.

Occasional confusion arises about ownership of certain pipeline tasks, especially during handoffs.

Never Rarely Sometimes Often Always Not sure

Some team members are aware of role definitions, but others are not, leading to inconsistent understanding.

Never Rarely Sometimes Often Always Not sure

Some teams follow the documented process, but others deviate or skip steps regularly.

Never Rarely Sometimes Often Always Not sure

PIPELINE FOUNDATIONS

Which of these exist today?

If it exists, name it and say where it lives. If it exists but is out of date or only partly done, mark it partial and say what is missing. If it is outside what you can check, say so rather than guessing.

Sales roles and responsibilities document

Fully in place Partial We do not have this Don't know — someone else would have to say

Team meeting attendance record

Fully in place Partial We do not have this Don't know — someone else would have to say

Training attendance records

Where the questions come from - and **what it asks of a respondent.**

AXAT questions are generated in a standardised format, so coverage is complete. You don't write them; you can add your own at the end.

HOW EACH CAPABILITY SECTION IS BUILT

Section intro & first question	YOURS TO EDIT	The framing, who should answer, and the perception question with its level options. Fixed per capability in the template.
Symptom blocks	GENERATED	Routed on the perception answer — respondents are only asked about the levels they claim and below. Volume follows the model's level count.
Evidence register	GENERATED	One question per artefact in the Progression Matrix: Fully in place · Partial · We do not have this · Don't know.
Closing open questions	GENERATED	Two free-text prompts per capability so respondents can add what the fixed questions missed.
Your own questions	OPTIONAL	Add extra questions at the end of the assessment, after the generated sections.

THE TRADE-OFF

STANDARD · QUICK

One perception question per capability.

A few minutes for any respondent. A working view, not a tested one.

AXAT · DETAILED AND DEFENDABLE

15 minutes to an hour per respondent.

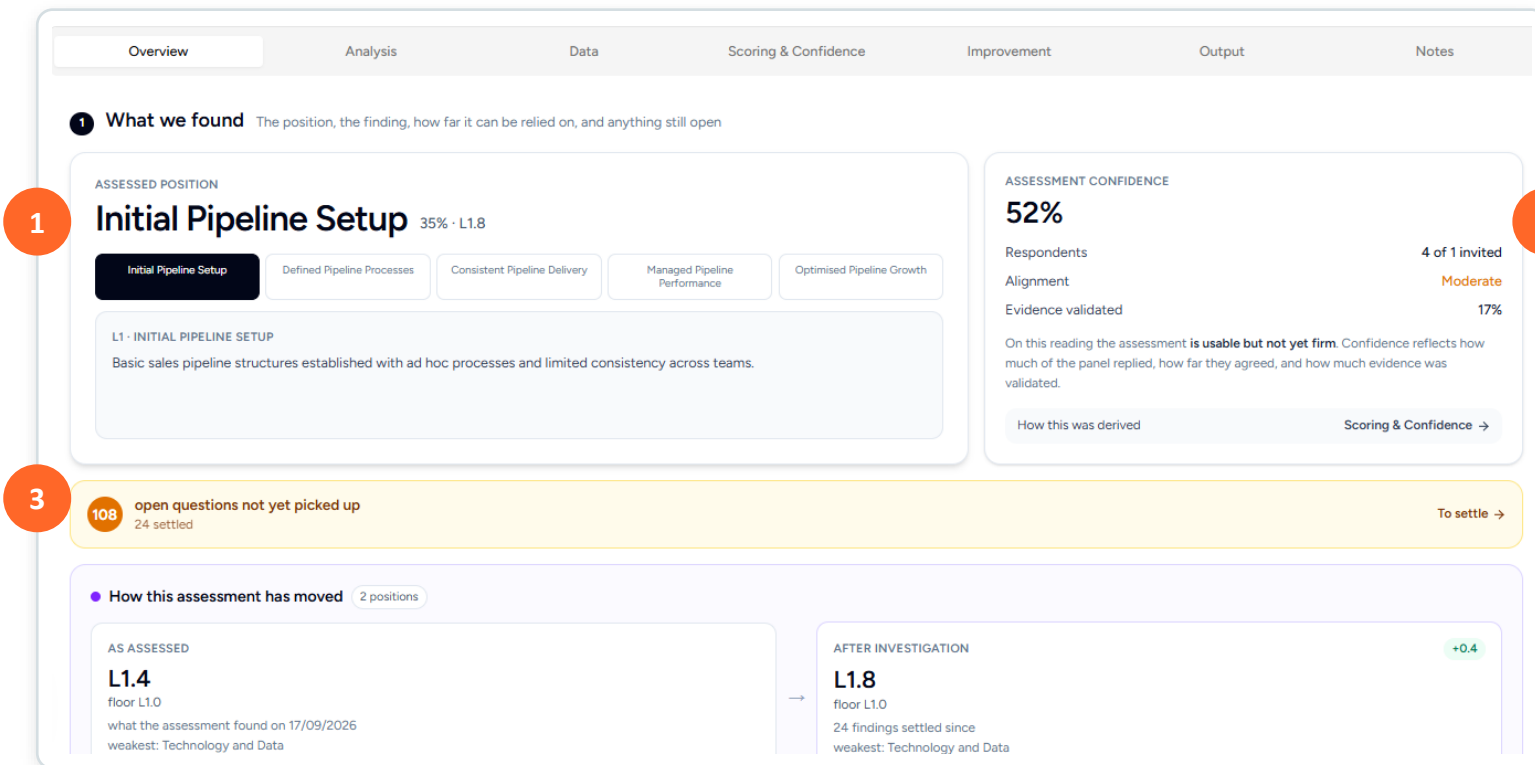
Depends on how many capability areas that person can cover, and the size of your Progression Matrix. In return: three lenses, a corroborated position and confidence you can defend.

Keep it manageable: scope the capabilities in the assessment, route respondents to the areas they know, and watch the question count in the wizard.

Standard when you need a quick read. AXAT when the finding has to survive challenge.

What you get back: a position, a confidence, and what's still open.

The Overview tab of a completed AXAT assessment. Everything the report will say starts here.



1. The assessed position

The corroborated level and the ladder it sits on, with the level description underneath - so the client sees what L1 means, not just the number.

2. Confidence, and what it's made of

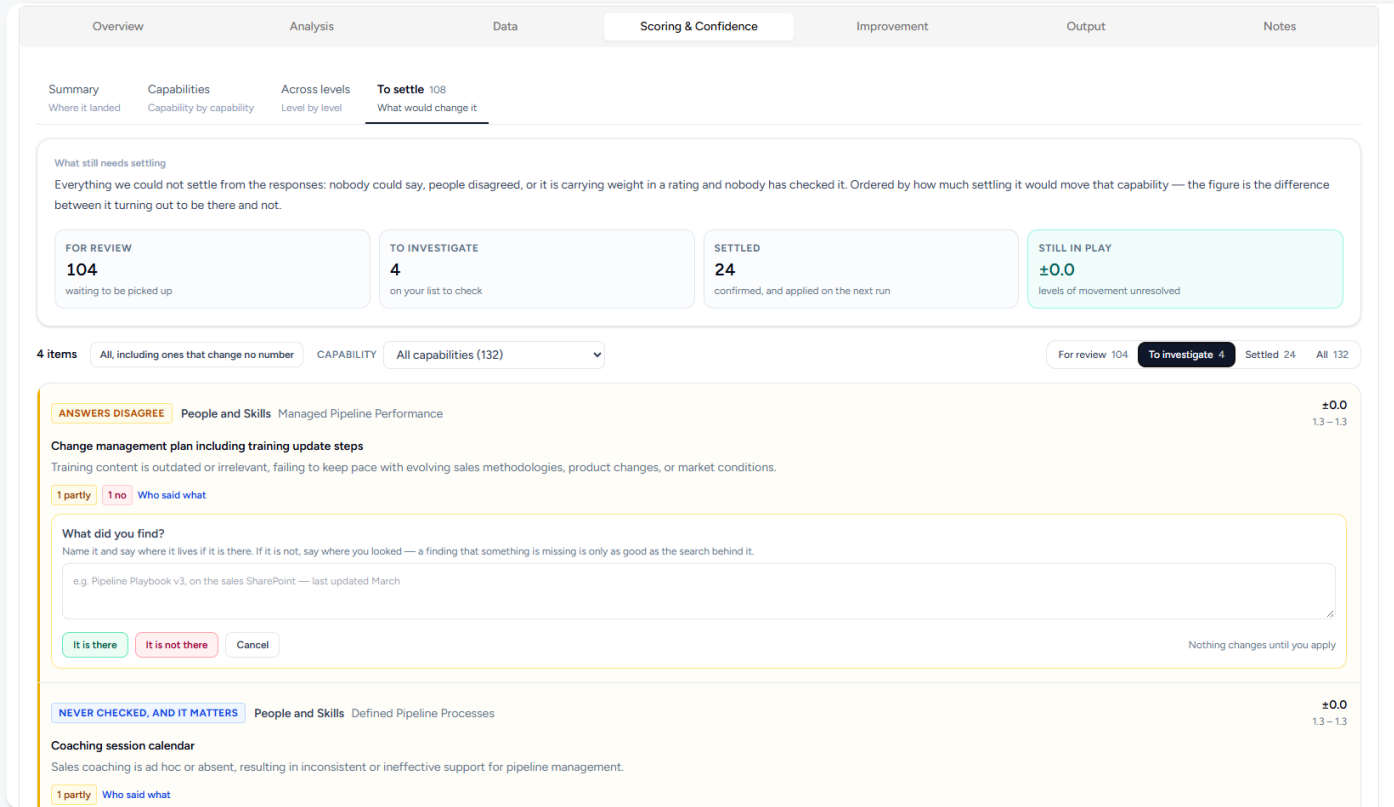
Respondents, alignment and evidence validated. "Usable but not yet firm" is the platform being honest about the reading.

3. Open questions, and the move

What's still to settle, and how the level shifted once findings were investigated - here from L1.4 as assessed to L1.8.

Scoring & Confidence: **check evidence** before it counts.

Respondents say an artefact exists. Before it sets a level, you check. The workflow lives inside the analysis view.



Overview Analysis Data **Scoring & Confidence** Improvement Output Notes

Summary Capabilities Across levels **To settle 108**

Where it landed Capability by capability Level by level What would change it

What still needs settling

Everything we could not settle from the responses: nobody could say, people disagreed, or it is carrying weight in a rating and nobody has checked it. Ordered by how much settling it would move that capability — the figure is the difference between it turning out to be there and not.

FOR REVIEW
104
waiting to be picked up

TO INVESTIGATE
4
on your list to check

SETTLED
24
confirmed, and applied on the next run

STILL IN PLAY
±0.0
levels of movement unresolved

4 items All, including ones that change no number CAPABILITY All capabilities (132)

For review 104 **To investigate 4** Settled 24 All 132

ANSWERS DISAGREE People and Skills Managed Pipeline Performance ±0.0
1.3 - 1.3

Change management plan including training update steps
Training content is outdated or irrelevant, failing to keep pace with evolving sales methodologies, product changes, or market conditions.

1 partly 1 no Who said what

What did you find?
Name it and say where it lives if it is there. If it is not, say where you looked — a finding that something is missing is only as good as the search behind it.
e.g. Pipeline Playbook v3, on the sales SharePoint — last updated March

It is there It is not there Cancel Nothing changes until you apply

NEVER CHECKED, AND IT MATTERS People and Skills Defined Pipeline Processes ±0.0
1.3 - 1.3

Coaching session calendar
Sales coaching is ad hoc or absent, resulting in inconsistent or ineffective support for pipeline management.

1 partly Who said what



Look into this

Flag an artefact that sets or blocks a level. It shows in the confidence reading until resolved.



Note

Record what you found - the draft was unpublished; the CRM fields exist only in the sandbox.



Confirm outcome

Present, partial or absent. Score and report update; the note carries into the findings as provenance.

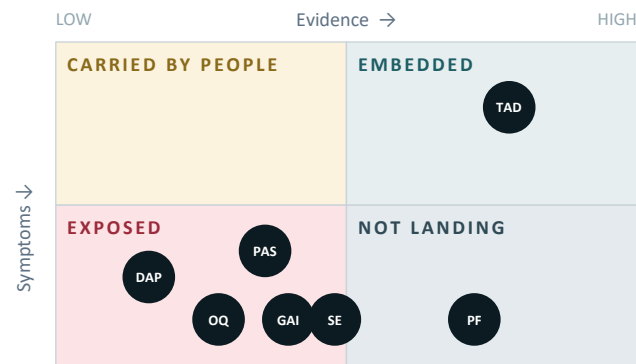
Validation is mandatory on the capability that sets the floor — the one an unverified artefact can move the most.

The score is the start. **Actionable advice** is the landing.

One assessment produces four things - and says what it does not yet know.

ILLUSTRATIVE · SUNSHINE POWER, A FICTITIOUS CLIENT

01 · WHAT KIND OF PROBLEM IT IS



Seven capabilities plotted by type of problem: documented evidence against lived experience.

Four kinds of problem need four different fixes. A maturity score can't tell them apart.

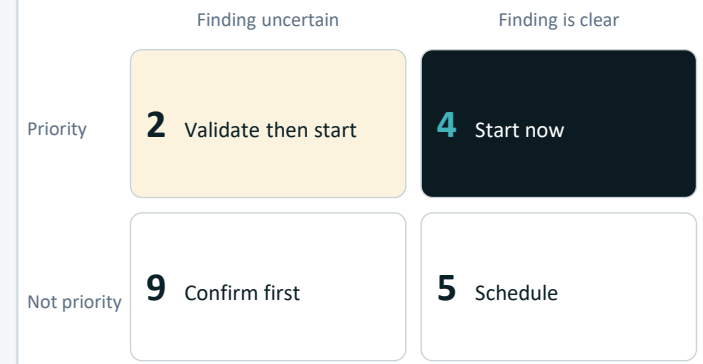
02 · THE ROOT CAUSES UNDERNEATH



73 named mechanisms found. 43 worth acting on.

Each failure mode is classified from the matrix, with the actions and artefacts that address it.

03 · WHERE TO START



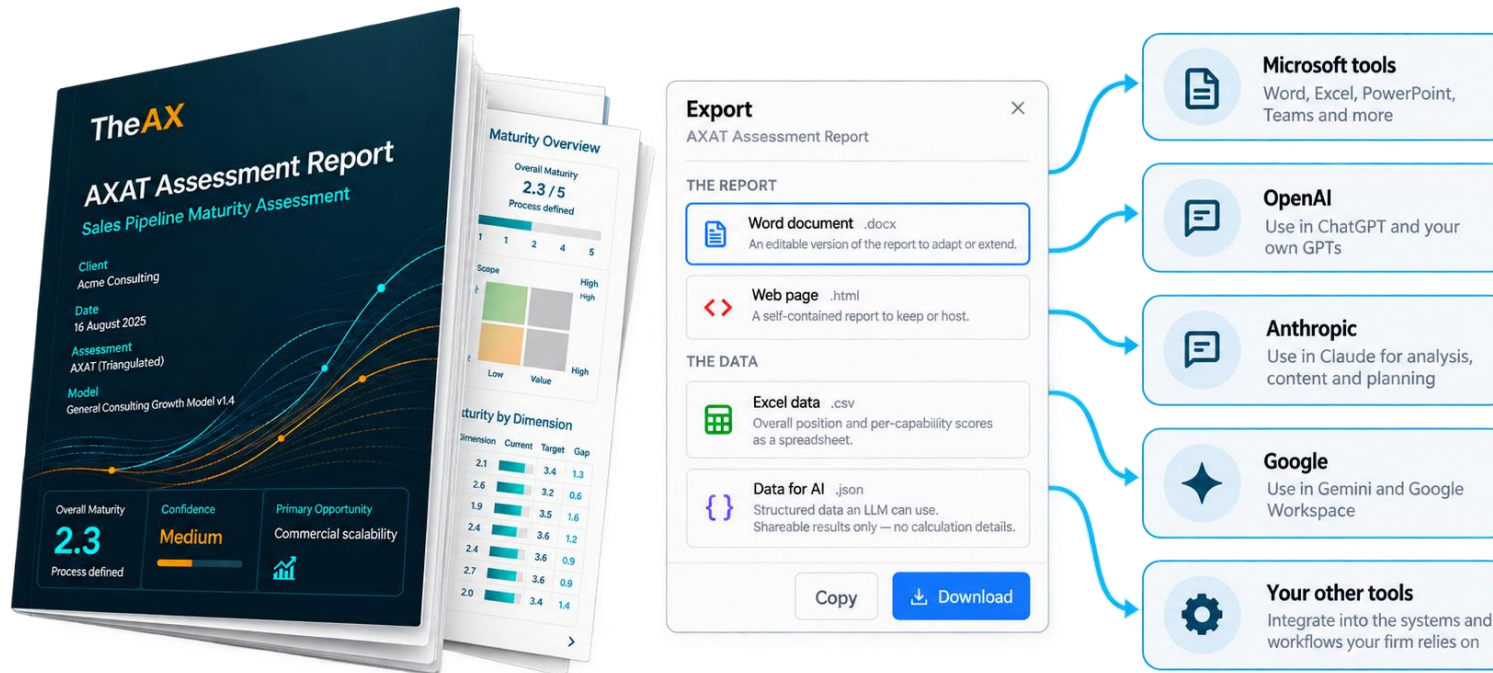
Twenty recommendations, sorted by priority and by whether the finding underneath is settled.

Clear guidance on what to act on first, and what to confirm before you do.

04 · AND WHAT IT DOES NOT KNOW: 42 of 73 questions open · 11 would move a rating · 30% of evidence independently validated

What comes out: a report the client can hold.

Generated from the analysis, so anything you settle in an investigation is reflected. Built from sections, branded once, shared or exported as a document or as data.



Fits your way of working
Export reports and data to use in any of your consultancy's workflows — including your preferred LLM.

WHAT'S IN THE REPORT

- **Summary and what the assessment found**
The corroborated position, the capability floor, the profile, and the readings behind each rating.
- **Where the wall sits, and what it means**
The next level each capability hasn't cleared; the Evidence x Symptoms state that says what kind of gap it is.
- **Root causes and capability summaries**
Every failure mode sorted by the change it asks for; a block per capability with its finding and actions.
- **Confidence and open questions**
What is settled, what remains open, what could still move a rating.
- **Actions, conclusion, appendix**
The recommendations selected; the recorded conclusion; every answer and result behind it, nobody named.

Build it from sections - core report, all sections, or any subset. Cover facts and branding are set once for every output. Share a link, export the report or the data, or import an amended copy back.

What to do first: **one book, one assessment, one report.**

Three steps, in this order. The first one is the only real work.

01

Author a Progression Matrix on one book

Pick the book you use most. Generate the matrix in the Workbench and validate the cells for the capabilities you assess most often. Readiness tells you when you're done.

02

Run one assessment in AXAT mode

Ideally internal or a friendly client first. Choose AXAT in the wizard, check the question count, send it to a handful of respondents.

03

Work Scoring & Confidence before the report

Investigate the artefacts that set the floor, confirm outcomes, then generate the report. Read the findings with the confidence panel beside you.

Launch page <https://theax.ai/releases/AXAT-Launch>

Contact us: Enquiries - hello@theax.ai. Customers - support@theax.ai