



AXAT

The AX Assessment Triangulation capability diagnostic

A diagnostic without evidence is an opinion.

AXAT is what makes it a measurement.

AXAT is the assessment engine inside TheAX, the assessment and diagnostic platform. It turns a consultancy's expertise into an automated, repeatable diagnostic that improves how consultants engage with clients. It **measures** capability through three independent lenses, **diagnoses** what is constraining the outcome, **advices** from the firm's own model, and **tracks** whether capability and business results move together.



MEASURE

Where the client actually stands

Capability read against your expertise, corroborated rather than self-reported.



DIAGNOSE

What is holding the outcome back

Not just a score - the cause, and which capability constrains the outcome.



ADVISE

The next best step

Guidance drawn from your firm's own model, specific to this client and their desired outcome.

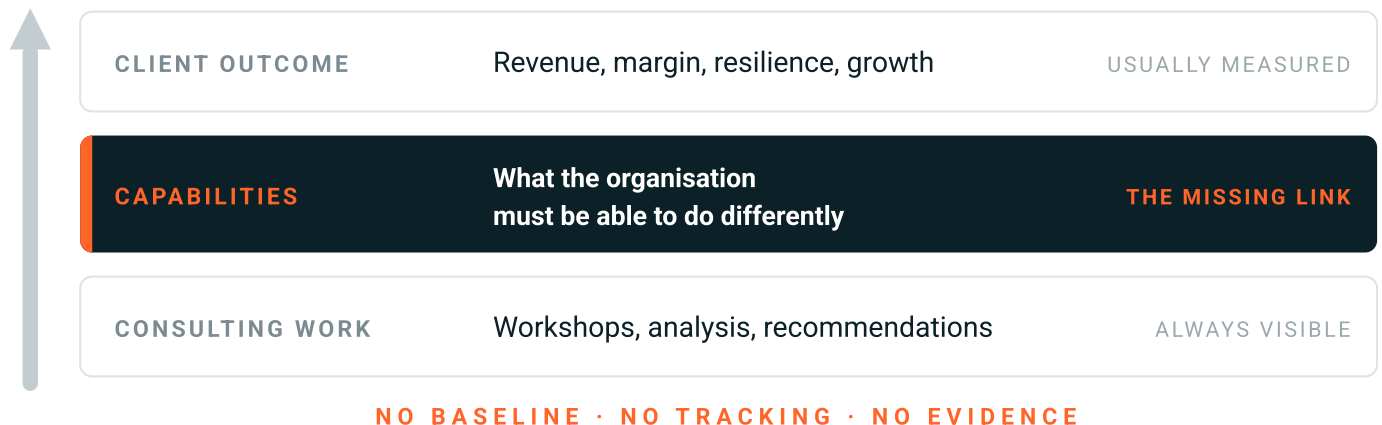


TRACK

Whether it worked

Re-read over time, so capability movement and business results are tracked and shown together.

01 · THE PROBLEM IT SOLVES



Capability is not one route to a client outcome among several. It is the only one — an outcome cannot change unless the organisation's ability to produce it changes. The work is visible. The outcome is visible. The layer that connects them is not measured at all.

Capability is never measured

No baseline at the start, no tracking through delivery, and no evidence the outcome is achieved.

The score cannot be defended

Measuring capability by asking is a claim, and a claim does not survive a challenge in front of a board.

The judgement cannot be delegated

Reasoning lives in senior people cannot be repeated reliably at volume, and it leaves when they do.

02 · HOW IT WORKS

Nothing can be automated while it lives in people's heads - that's why you create Playbooks. But for consultancy to be scalable expertise has to be written down in a form a machine can reason over. Everything after that is generated from it.

STAGE 01



The Consulting Atomic Model

THE BASE

An LLM by itself cannot reason with your expertise and make it repeatable and measurable. That can only be achieved when your expertise is structured. The **Consulting Atomic Model** is that structure: capabilities and levels, the evidence that proves them, the guidance that improves them, and the pathways connecting capability to outcome. It is curated and versioned in your firm's Product Book. TheAX uses AI workflows to structure and store your expertise in the way that makes AXAT diagnostics possible. Authored once with your senior people, in your language. **Everything downstream is generated from this** — which is why the answer is your firm's judgement rather than a public model's. This is your IP that improves with every engagement.

STAGE 02



The assessment

GENERATED, NOT WRITTEN

To make information collection repeatable and automated TheAX builds a question-based assessment from your model: Perception scoring per capability, the artefacts to request, and the symptom probes that catch what documents miss. Send to stakeholders as an automated interviewer or consultants to use the assessment in face-to-face interviews. **Nobody writes a bespoke questionnaire per client**, and every engagement asks in the same disciplined way.

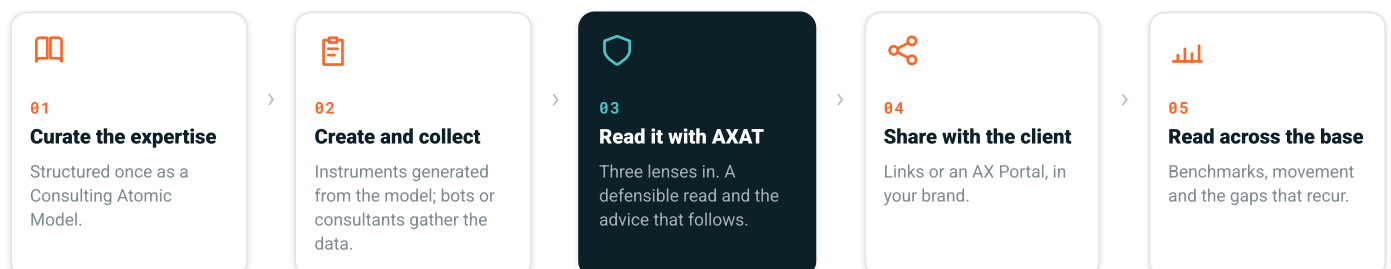
STAGE 03



AXAT reads it

THE ALGORITHM

Consultants have more important work to do with clients than spending days crunching numbers and performing diagnostics that AXAT makes repeatable and deliverable instantly. The algorithm triangulates the three lenses to create meaningful and deep diagnostics that help consultants get the punchline faster, and with diagnostics that can stand-up to client scrutiny. It runs on every capability, for every client, at machine speed — and reports how confident it is in each reading.



One curated model runs the assessment to collect the data, produces the outputs, and stays behind the client-facing AI. **The consultancy's judgement is present at every step — including the ones it is not in the room for. So juniors to seniors can run these assessments with confidence**

03 · THE TWO SCORING RULES

RULE 1 — CAPABILITY LEVEL

The weaker lens wins.

A capability scores at **the lower of evidence and symptoms**, never at the claim. Evidence sets the floor; symptoms cap it. The claim is held separate and tested against both.

RULE 2 — FIRM LEVEL

Levels are cleared, not averaged.

A level is reached only when **every capability has reached it**. Strength in one area does not buy back weakness in another. A mean is used as context, not as the answer.

What it asks of you. AXAT costs no more to run than a self-rating survey, and reaches a meaningful diagnostic faster. It asks a little more only where you choose to close the evidence gaps it identifies, which is what lifts the confidence score — and that is the point, because the evidence layer is where the value sits. The effort is front-loaded into building your model: once that exists, every new assessment is configured in minutes.

04 · WHY THREE LENSES

A score is only as strong as the evidence beneath it.

A doctor never diagnoses from one answer.

They ask how you feel, examine you, then send you for tests — three instruments, none trusted alone. A good consultant already works this way. **What they cannot do is work that way on every capability, for every client, at speed.** That is the part AXAT automates.



WHAT PEOPLE SAY

Perception

Leaders and teams rate the capability. The score **and the spread of opinion** stay visible instead of being averaged away.



WHAT CAN BE SHOWN

Evidence

Named artefacts, marked in place, partial or absent. Not "do you have a process?" but **"show me the one you used last month"**.



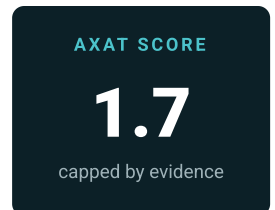
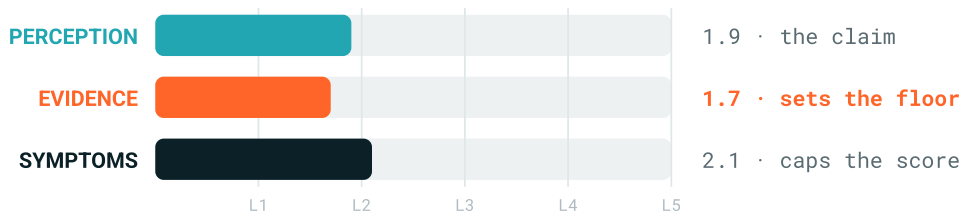
WHAT IS EXPERIENCED

Symptoms

If this capability were weak, what pain would surface? Enough stakeholders reporting that pain **counts as a signal**, whatever the documents say.

ONE CAPABILITY, THREE READINGS

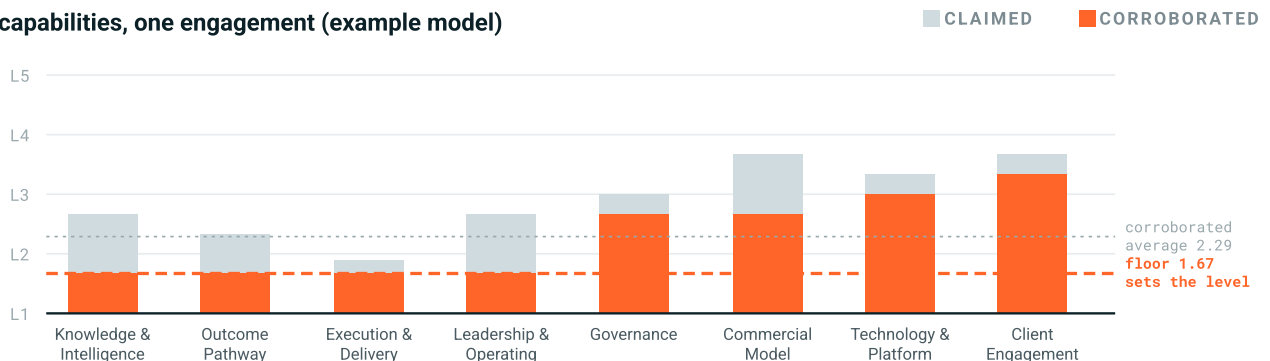
Execution & Delivery



The stakeholders claims L1.9 and symptoms suggest L2.1, but the artefacts only support L1.7 — so the capability reads at **1.7**. Had the artefacts been strong and the symptoms weak, symptoms would have capped it instead. **Disagreement is reported as a finding, not averaged away.**

05 · WHAT THE RULES DO TO A REAL NUMBER

Eight capabilities, one engagement (example model)



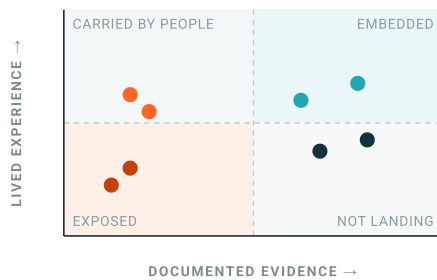
PANEL'S OWN AVERAGE **L2.90** | CORROBORATED AVERAGE **L2.29** | LEVEL ACTUALLY CLEARED **L1**

Four capabilities sit on the floor at 1.67, so nothing has cleared L2 and the firm reports at L1. Averaged instead, the same engagement returns L2.29 — more than a level above the position that exists; the panel's own average, L2.90, is nearly two. Anonymised engagement, consultancy growth model.

06 · DIAGNOSTIC INSIGHT

Four kinds of gap. Four different fixes.

AXAT goes beyond a score. The AI reads across the assessment to surface cross-cutting themes and identify the root cause of what is holding each capability back. Documented evidence plotted against lived experience sorts every capability into one of four problems, each needing a different intervention. **Two of the four are not fixed by writing anything down.**



● Carried by people

No system behind it. Delivery feels acceptable because it lives in individuals, and works until those individuals are busy, absent or gone.

● Exposed

No system, and delivery already shows it. Symptoms are visible to the client now, not later.

● Embedded

System and practice aligned, and ahead of the rest of the firm. Protect these — do not spend scarce change capacity here.

● Not landing

Artefacts exist and authoring is done, but embedding has not happened. This is an adoption problem.

07 · WHAT YOU GET OUT

Not only faster consulting. A more scalable firm.



A position you can defend

Every capability scored, with the lens that set the level named and confidence reported separately.



A diagnosis, not a score

Which capability is constraining the outcome, and which kind of gap it is, so the fix matches the cause.



Repeatable advice

The next best step from your own model, so a mid-weight team gives the answer your best partner would.



Benchmarking that means something

Same model, same rules every time, so readings compare across clients, business units, suppliers and time.



Outputs in usable formats

Reports and board packs in your brand, produced from the assessment — plus exports for other workflows and LLMs.



Sharing that does work for you

Share document links or an AX Portal where the client explores the data and asks an AI grounded in your expertise.

More on TheAX

Every part of the platform has its own fact sheet, and a walkthrough on a completed assessment is the quickest way to judge the method for yourself. Both start at the website.

hello@theax.ai · Marlborough, United Kingdom

FIND EVERYTHING AT

theax.ai