

AI Operating Impact Briefing

This Briefing reads your operation through Envisago's proprietary AIVOM™, the AI Value Operating Model. AIVOM™ examines four connected dimensions that together determine whether AI produces operating impact: Value, Operating Design, Capability, and Performance. Strength in one rarely delivers on its own, because the dimensions reinforce or constrain one another. The profile below shows where your operation stands across each dimension, and the Briefing explains how they connect.

SAMPLE BRIEFING · A completed Briefing for a fictional operation: the customer operations function of a mid-market insurer. Nothing here describes a real company. Your Briefing reads your own operation, from your own answers.

Your Operation

This briefing is for a customer operations function in insurance, handling policy servicing, claims enquiries, and renewals across phone, email, and chat for retail customers. At 100 to 250 people processing up to 50,000 contacts a month, with 12 to 18 months of deliberate AI deployment behind it and a mixed relationship with AI across the team, this is an operation that has moved beyond early experimentation but is not yet seeing the operating impact the investment was intended to produce.

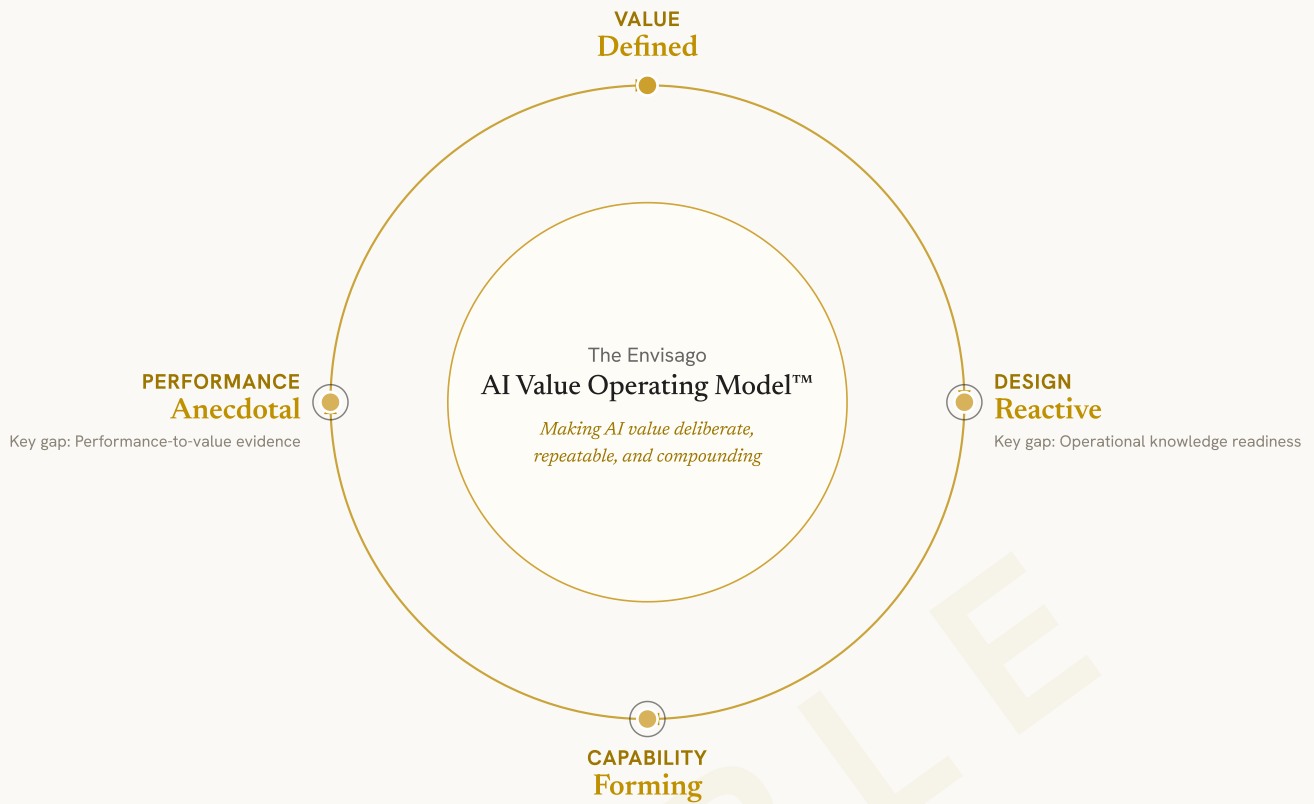
Your AIVOM™ Impact Profile

AIVOM™ LOOP STATE

Emerging

Fragmented · **Emerging** · Connecting · Compounding

The foundations of the loop are beginning to form, but the connection between AI activity and operating impact is still uneven.



YOUR POSITION ACROSS THE FOUR AIVOM™ DIMENSIONS

Value clarity

Defined

Unclear · Directional · **Defined** · Aligned

AI is connected to specific objectives and the value picture is taking shape, though it is not yet complete or consistently tied to goals across the operation.

Operating design

Reactive

Unadapted · **Reactive** · Redesigned · AI-native

AI is still largely layered on top of the operation, so the supporting conditions for AI-enabled work are adapting reactively rather than by deliberate design.

Capability readiness

Forming

At risk · **Forming** · Enabled · Distributed

Willingness is growing and AI capability is developing in pockets, but it is not yet spread or supported systematically.

Performance evidence

Anecdotal

Unproven · **Anecdotal** · Measured · Proven

Some evidence exists, but it remains partial, uneven, or not yet strong enough to give a clear view of impact.

WHAT IS SITTING BENEATH THE LOOP

Below is what sits beneath each part of your AIVOM™ profile. The bars show where each answer landed across the twenty sub-dimensions. **Highlighted areas are the gaps selected as most material to this Briefing.**

Each row uses a four-step scale, from early-stage on the left to advanced on the right.

How to read this section

Your three deepest gaps are highlighted below. Where several score low, priority goes to the areas inside the dimension holding back your loop as a whole. Low scores across dimensions surface in how the gaps connect.

VALUE	Defined	OPERATING DESIGN	Reactive
Strategic alignment		Workflow design	
Value goals		Data and systems readiness	
Investment and cost visibility		Operational knowledge readiness	
Accountability		Quality approach	
Value evidence		Governance	
CAPABILITY	Forming	PERFORMANCE	Anecdotal
Culture and readiness		Financial performance	
Training and development		Operational performance tracking	
Capability distribution		Experience tracking	
Role redesign		Innovation and new value	
Team Performance		Performance-to-value evidence	

WHAT IS HOLDING THE LOOP BACK

Operating design, especially operational knowledge readiness

The operation may know what value it wants and have capable people, but it is not yet built to deliver that value reliably with AI. Until workflows, data, operational knowledge, quality, and governance are designed for AI rather than layered on top, intent and capability do not convert into dependable performance.

Across the four AIVOM™ dimensions, Value clarity is the strongest position here. Strategic alignment is in place, value goals have been set, and there is senior ownership with accountability distributed into the operation. That is a meaningful foundation, and it is more than many operations at this stage have established.

The difficulty is that this strength is not converting into operating impact, because the dimensions that should carry it forward are weaker. Operating design sits at Reactive, and the driver is not workflow structure or governance, both of which are more developed, but operational knowledge readiness, which is the foundation the rest of the design depends on. Capability readiness is Forming, held back by concentrated AI capability and performance expectations that have not yet caught up with how work has changed. Performance evidence is Anecdotal: the operation reports handling times and CSAT but cannot connect either to what AI is specifically contributing. The binding constraint is Operating design, and within it, the state of operational knowledge is the single factor most limiting what the rest of the operation can do. Until the knowledge AI relies on is explicit, current, and AI-readable, the assistant pilot will continue to produce variable results regardless of how much is invested in the layers above it.

Your Highest-Impact Gaps

Operational knowledge readiness

The generative AI assistant pilot is producing variable results depending on who is using it, and the most likely explanation sits in what the tool is operating on rather than in the tool itself. When operational knowledge is patchy, out of date, or written for experienced humans rather than for AI, the assistant fills the gaps with its own judgement. Handlers who know the business well can catch and correct that. Those who are less experienced, or working on less familiar policy types, cannot. The result is inconsistency that looks like a capability problem but is actually a knowledge problem.

At 20,000 to 50,000 contacts a month, this is not a marginal issue. Each interaction the assistant handles is drawing on the same incomplete knowledge base, and the errors it produces at that volume are systematic rather than random. Progressing here means auditing what the assistant actually relies on: process guides, decision logic, escalation rules, product-specific handling instructions. For each, the question is whether it is current, whether it is explicit enough for AI to apply without human interpretation, and where the gaps are. This does not require a complete overhaul before anything else moves. It requires knowing which knowledge gaps are producing the most inconsistency and addressing those first.

Capability distribution

Two people are carrying the operation's AI capability, and the rest of the team routes through them. That is a known risk, and the team knows it. The concern is not just that those two individuals could leave, though that is real. It is that at this scale, two people cannot be the decision point for AI-related questions across 100 to 250 handlers without becoming a bottleneck that slows adoption and concentrates risk.

The quiet anxiety some of the team feel about what AI means for their roles is also relevant here. Where people are uncertain about whether engaging with AI makes them more or less secure, they engage less, share less, and defer more. That is a rational response to an unresolved question, not a capability failure. Progressing on distribution means two things running in parallel. First, making AI interactions visible: the two people who are effective with the assistant should be sharing how they work with it, not just the outputs. Second, the operation needs to give the wider team a clear answer to the role question, a specific account, rather than a reassurance, of how roles are changing and what the operation expects from people working alongside AI.

Performance-to-value evidence

The board is asking what the operation has to show for the spend, and the honest answer at present is that it cannot demonstrate it. Handling times and CSAT are reported monthly, but neither connects to what AI is specifically contributing. The operation has value goals, but there is no mapping from those goals to the data that would confirm whether they are being met.

This is not a measurement failure in isolation. It reflects the sequence in which things were built: value goals were set, tools were deployed, and the evidence infrastructure was not established alongside them. Progressing here starts with a straightforward mapping exercise: for each value goal the operation set at the outset, what would evidence of achievement look like, and is that data available from existing systems? Some will be reportable now with modest effort. Others will require system or reporting changes. The operation should know which is which, so it can produce credible evidence for the goals it can already demonstrate and have a clear plan for the rest.

Cross-Dimensional Insight

The operation defined value goals at the outset, which is the right starting point. But those goals have not been connected to any performance measure that would confirm whether they are being achieved. This matters beyond the board conversation. Without that connection, the operation cannot tell which parts of the AI deployment are working and which are not, which means decisions about where to invest next, what to scale, and what to change are being made without the evidence to support them. The goals exist and the performance data exists; the link between them is what is missing.

Where to Focus First

Given the operation's scale and the gaps identified, the highest-impact starting point is the operational knowledge base, specifically identifying which knowledge gaps are producing the most inconsistency in the assistant pilot and making those explicit and AI-readable.

The reason this takes precedence over the evidence gap, which is the more visible pressure given the board's question, is that the knowledge problem is upstream of almost everything else. Better evidence will show the operation more clearly what is not working. But if the knowledge base remains patchy, fixing the evidence layer will surface a problem the operation is not yet positioned to solve. Addressing knowledge first means that as capability distributes and evidence improves, the assistant is operating on a more reliable foundation, and the results being measured are closer to what the tool can actually deliver.

The constraint on taking people offline is real, and this does not require it. The most effective approach at this stage is pairing the two people who are already effective with the assistant with a structured brief: document the knowledge gaps they work around every day, the places where they add context the assistant does not have, and the decisions they make that the tool cannot. That is the audit, and it can happen alongside normal operations. The change fatigue after the platform migration is also worth respecting: this is a targeted intervention on the specific thing most limiting what the operation already has in place, rather than a transformation programme.

NEXT STEP

Go deeper with the AI Operating Impact Roadmap: a structured AIVOM™ deep dive into the gaps that matter most for your operation, how they connect, and the right order to address them, with practical first steps for turning AI activity into operating impact.

[Get the full Roadmap](#)

This Briefing was generated by AI using Envisago's proprietary AIVOM™ methodology.
AI can make mistakes, so discretion is advised.