

INSIGHT PAPER 03

Data readiness

Your data can be clean, governed and warehoused, and still be unable to answer the question an AI agent needs answered at the moment a customer asks it.

Rachel Lane, Founder and CEO, CXais

EXECUTIVE SUMMARY

Four statements

The problem. Data quality is not data readiness. A quality audit asks whether the data is correct. AI deployment depends on something else, whether the data can evidence a specific decision, for a specific customer, at the moment of contact.

The finding. Readiness is not an enterprise condition. It has to be assessed workflow by workflow, because the same field passes for one workflow and fails for another.

The framework. Five tests decide it. Reachability, recency, identity, authority, permission. A workflow is only as data-ready as its weakest required test.

The implication. Most readiness failures do not call for enterprise-wide data transformation. They call for targeted remediation around the workflow being automated, usually measured in weeks.

The cost of not testing is that the remediation happens anyway. Later, in public, with a customer waiting.



The question the audit does not ask

Every enterprise AI programme runs a data readiness exercise. Almost all of them run the wrong one.

The exercise that gets run is a data quality audit. Completeness, accuracy, duplication, lineage, governance sign-off. It is usually owned by the data function, it produces a scorecard, and the scorecard usually comes back acceptable. The programme proceeds on the strength of it.

Then the agent goes live and cannot tell a customer whether their broadband will be working on Friday.

Nothing in the audit was wrong. The audit answered the question it was asked. The problem is that data quality and data readiness are different tests, and only one of them predicts whether AI will perform in a customer conversation.

Two instruments, two questions

A data quality audit asks whether the data is correct. A data readiness assessment asks whether the data can be used to make a specific decision, for a specific customer, inside a specific time budget, with a defensible basis for the answer. Those are not variations of the same question.

What the scorecard cannot see

A readiness assessment asks four things a quality scorecard never records. Can this be retrieved while the customer is still on the line. Is it true at the moment the answer is given rather than at the moment it was extracted. Which system wins when two of them disagree. And is this use of it permitted here. None of those are properties of the data. They are properties of the data in use, which is why they only appear when a workflow is named.

Why the quality audit passes

The audit is not failing. It is answering correctly for the workload it was designed around. Reporting and analytics are batch, aggregate and retrospective. They tolerate overnight refreshes, they work on populations rather than individuals, and nothing breaks if a source is reconciled tomorrow rather than now.

A customer conversation is the opposite on all three counts. It is runtime, it concerns one individual, and it cannot wait for reconciliation. Data that is entirely fit for the first

workload can be unusable in the second, and no measure of completeness or accuracy will reveal that, because completeness and accuracy are not the properties under strain.

EXHIBIT 1

DATA QUALITY AUDIT	DATA READINESS ASSESSMENT
Is the data correct, complete and governed?	Can this workflow's decision be evidenced at the moment of contact?
Unit of analysis is the system or the domain	Unit of analysis is the workflow
Owned by the data or technology function	Owned by the programme, with data and operations together
Produces an estate-wide scorecard	Produces a verdict per workflow, with a named remediation
Misses latency, source conflicts, permission scope and the interaction estate	Says nothing about long-term data governance, which still matters
Right instrument for platform investment, regulatory reporting, migration	Right instrument for AI deployment decisions

Both are legitimate instruments. Only one of them tells you whether to deploy.

Data readiness for AI in customer experience is not an enterprise condition. It is a per workflow condition.

That distinction is what connects this paper to the last one. Until you know which workflow the AI is performing, which decision sits inside it and which action follows, there is nothing specific enough to test. Workflow cartography comes first because data readiness has no unit of analysis without it.

THE FRAMEWORK

The five tests of AI data readiness

Reachability. Recency. Identity. Authority. Permission. Every workflow, tested before it is automated.



The five tests of AI data readiness

Data readiness for a workflow resolves into five tests. Each returns a pass, a conditional pass or a fail.

1

Reachability

Can the AI retrieve the data it needs, in the workflow, at the moment it needs it?

Not does the data exist. Not is it in the warehouse. Can a system call it, get an answer back, and use it before the customer disengages.

What fails it. Data that lives only in a screen a human reads. An API rate-limited for a nightly job. A system whose owner will not grant service account access. A response that returns in eight seconds when the interaction allows two. Data held only as an aggregate when the decision needs the individual record.

2

Recency

Is it current enough for the decision being made?

Every source has a staleness profile, and whether that profile is acceptable depends entirely on what is being decided. Yesterday's address is usually fine. Yesterday's appointment availability is not. This is where the decision time budget belongs, set by the consequence of being wrong and the speed at which the underlying state changes. **Data readiness has no universal freshness standard. The workflow sets the time budget.**

What fails it. The agent gives an answer that was true this morning. The customer acts on it. The organisation has now told a customer something untrue, in writing, at scale, with a timestamp on it.

3

Identity

Can the system establish that this is the same customer, across the channels the workflow touches?

The caller is known by their calling line identity. The web session is known by a cookie and a login. The email arrives from an address that belongs to their partner. The chat is anonymous until authentication, which happens three exchanges in.

What fails it. Weak resolution produces one of two outcomes. The agent asks for information the customer has already given, which destroys the experience the programme was bought to improve. Or it joins two customers together, which is a data incident with a reporting obligation attached.

4

Authority

When two sources disagree, which one is trusted, and is that rule written where a machine can read it?

Every enterprise has facts that live in more than one place. Address in the CRM and in billing. Contract end date in the quoting tool and in the service platform. Consent state in the marketing platform and in the service record. Humans resolve these conflicts constantly and invisibly. Nobody has ever asked an advisor to explain it, so nobody has written it down.

What fails it. An AI agent has no such judgement. It takes whichever source it was pointed at, or it takes both and contradicts itself inside the same conversation. This is the test mature estates fail most, because more systems means more places for the same fact to live.

5

Permission

Is the AI authorised to access and use this data, for this purpose, in this channel, for this customer?

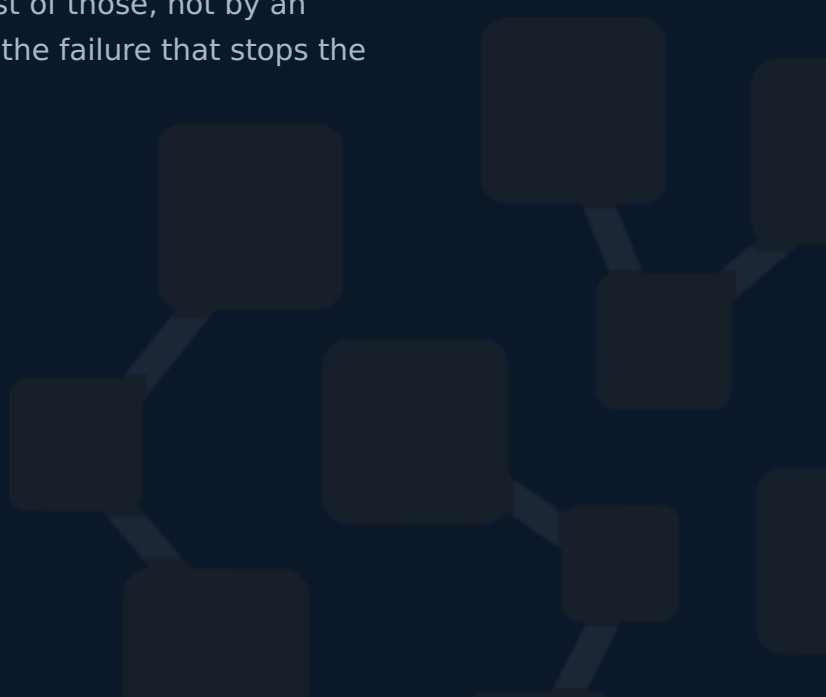
Technically accessible is not the same as operationally, legally or ethically permissible. Consent, purpose limitation, retention and special category handling attach to specific uses, not to the estate as a whole. Data lawfully held for service delivery is not automatically available for a proactive outbound conversation.

What fails it. The consent flag lives in the marketing platform. The workflow runs in the service platform. The agent cannot see the flag it is required to honour.

THE PRINCIPLE

A workflow is only as data-ready as its weakest required test.

The word required matters. Not every test binds on every workflow. A post-authentication account query carries no identity risk. A workflow reading a single system of record has no authority conflict to resolve. Testing tells you which of the five apply, and readiness is set by the worst of those, not by an average across all five. Averages hide the failure that stops the workflow.



SECTION 03

One customer, one question, four systems

Return to the broadband question. Four systems sit behind it, and the customer asking it has no idea.

A telecoms operator whose customer record is 99.4% complete, and which passes every governance test in the estate. An AI agent is asked, on a Tuesday afternoon, whether the customer's broadband will be live on Friday.

To answer, the agent needs the order status from the CRM, the provisioning state from the operational support system, the engineer appointment from field scheduling, and any open exception on the line from the port tracker. Four systems, one apparently simple question, asked thousands of times a week.

EXHIBIT 2

REQUIRED DATA	REACHABILITY	RECENCY	IDENTITY	AUTHORITY	PERMISSION
CRM order status	Pass	Pass	Pass	Unresolved No precedence rule against the OSS	Pass
Provisioning status	Pass	Fail Overnight refresh only	Pass	Unresolved Disagrees with CRM order state	Pass
Engineer appointment	Pass	Conditional Hourly refresh against a two hour promise	Pass	Pass	Pass
Port exception	Fail No callable interface, screen only	Pass	Pass	Pass	Conditional Third party data, purpose scope unconfirmed

Four sources, five tests, twenty answers. Fourteen are clean. The workflow still cannot run, because readiness is set by the six that are not.

THE VERDICT

DATA QUALITY

99.4%

Complete, governed, signed off by the data function.

AI DATA READINESS

Fail

Two failed tests and two unresolved authority questions.

The same workflow, measured two ways, on the same day. One of these numbers went to the steering committee. The other one decided whether the deployment worked.

Note what the verdict does not say. It does not say the workflow cannot be automated. It says three specific things have to happen first, and each of them is a piece of work with an owner and a duration.

EXHIBIT 2B, FROM VERDICT TO SCOPE

WHAT FAILED	WHAT HAS TO HAPPEN	WHO OWNS IT	EFFORT
Port exception (reachability)	Expose the port tracker through a callable interface, or route the exception check to a human step until it exists	Integration, with network operations	Weeks
Provisioning status (recency)	Move the refresh inside the promise window, or design the conversation to report status without committing to a date	OSS owner, with conversation design	Days to weeks
CRM against OSS (authority)	Decide which system wins on order state, write the precedence rule with a divergence window, configure it	Operations lead, with data	Days

None of these is a data platform programme. Two of the three are decisions somebody already makes informally, written down for the first time.

That is a scope. Without the assessment, it is an incident.

What we see in the field

Three patterns recur across our assessments, consistently enough to plan around.

The failure is rarely accuracy

It is reachability or authority. The data exists, it is correct, and it either cannot be retrieved inside the time the conversation allows, or it exists in two systems that disagree with no rule for which one wins.

The interaction estate is the least ready and the most assumed

Organisations that would never deploy against unvalidated CRM data will point an AI system at ten years of unlabelled call transcripts and expect insight from it.

Readiness is the dimension most likely to be deferred

Remediation looks like infrastructure work, and infrastructure work is slow. Deferring it does not remove it. It relocates it, from a pre build workstream with a plan to a post go live incident with an audience.

THE PUBLISHED EVIDENCE IS CONSISTENT WITH ALL THREE

95%

of enterprise generative AI pilots delivered no measurable P&L impact

MIT Project NANDA, 2025

80%+

AI project failure rate, roughly twice that of non-AI IT projects

RAND Corporation, 2024

42%

of companies abandoned most of their AI initiatives in 2025, up from 17% the year before

S&P Global Market Intelligence, 2025

2x

more likely to have redesigned workflows before selecting a modelling approach, among organisations reporting significant returns

McKinsey, 2025

None of those studies measured model capability. They measured what happened around the model.

The interaction estate

Everything so far concerns structured data. The larger and less examined half of the CX estate is the interaction record itself.

Calls, chats, emails, messages, advisor notes, survey verbatims. Organisations assume this estate is an asset because it is large. Volume is not readiness. Three problems recur.

Outcomes are not labelled

The recordings exist. What happened as a result of them mostly does not. Whether the issue was actually resolved, whether the customer came back about the same thing nine days later, whether the promised action was completed, is either absent or sits in a separate system with no reliable join. An estate with no outcome labels can describe what was said. It cannot be used to learn what works, which is the thing it was assumed to be for.

Metadata is thin or wrong

Disposition codes are chosen by advisors under handling time pressure, which makes them a record of advisor behaviour as much as customer intent. Where one code carries a disproportionate share of a queue, it is usually functioning as the default when the right code is unclear, and every analysis built on it inherits that.

Coverage is not what the estate believes

Quality assurance sampling is not a representative sample of contacts. It is a sample of contacts selected for review, historically weighted towards particular teams, shifts and outcomes. Training or evaluating an AI system against it inherits that weighting silently.

There is a timing argument too. IBM's Institute for Business Value puts the average useful lifecycle of an AI model at approximately 14 months. An interaction estate that takes two years to label and structure is being prepared for a model generation that will be retired before the work completes. That argues for narrow, workflow-specific labelling ahead of general estate remediation, which is where the five tests point as well.

THE REMEDIATION RULE

Fail does not mean stop.

A workflow that fails one test is a workflow with a date, once somebody owns the remediation. Treating a failed test as a verdict on the use case is how organisations end up with a shortlist of low-value workflows that happened to have tidy data.



From readiness gap to remediation

A failed test is not a two year data transformation. It is a named piece of work, and naming it is most of the value.

The cheapest remediation available is writing precedence rules. Every authority conflict you find is a rule that already exists in somebody's head. An advisor knows which system to believe. Getting that into a document, and then into configuration, costs days. It almost never appears on a data programme plan, because it is not a technology task and no technology function owns it.

Three rules for writing precedence rules

Precedence is set per field, not per system. Billing can be authoritative for address and wrong about contact preference, at the same time.

Every rule needs a divergence window. Two sources that reconcile within an hour need a different rule from two that can disagree for two days.

Where no rule can be agreed, the workflow does not consume that field. Saying so explicitly is a design decision. Leaving it unresolved is a defect waiting for a customer to find.

A worked instance, from practice

A payment dispute workflow I worked on at a retail bank would have passed reachability, recency, identity and permission, and failed on authority alone. Dispute status existed in both the case management system and the card scheme portal, and the two could disagree for more than a day with no precedence rule between them. The remediation is a rule and a configuration change. Weeks, not quarters. Left unfound, a workflow like that either goes live contradicting the customer's own app, or sits behind a platform initiative that nobody sequenced against it.

Data readiness across the maturity model

The five canonical maturity levels apply to this dimension as they do to the other five.

Most organisations we assess sit at Explore or early Build here while describing themselves as further ahead, because their platform maturity is genuinely high. Platform maturity and workflow readiness are different measurements.

01	Explore	Data quality is discussed as an estate-wide condition. No workflow has been tested for runtime retrieval. Interaction data is retained, not structured.
02	Build	Individual workflows have been tested. Reachability gaps are known and being closed for a first wave. Authority conflicts are documented and precedence rules drafted.
03	Operational	Priority workflows pass all required tests. Runtime retrieval is monitored. Consent state is visible in the channel where the decision is made.
04	Optimised	Outcome labelling feeds retraining for deployed workflows. Time budgets are defined per workflow and tested against source refresh rates.
05	Leading	Readiness is assessed before any workflow enters the automation roadmap. New sources arrive with a precedence rule and permission scope defined.

THE SEQUENCING RULE

Value at stake $\times$
readiness $=$ roadmap

One workflow is a decision. Forty workflows is a roadmap, and it is built from two numbers, not one.



SECTION 08

Readiness across the portfolio

Read across a portfolio, the five tests answer the question an executive actually has.

EXHIBIT 3

WORKFLOW	WEAKEST TEST	READINESS	REMEDICATION	VALUE
Order status enquiry	Reachability	Not ready	Callable interface for port exceptions	High
Payment dispute	Authority	Not ready	Precedence rule and configuration	High
Appointment rebooking	Recency	Ready with constraint	Conversation designed to avoid a two hour commitment	Medium
Billing explanation	None binding	Ready	None	Medium
Proactive retention outreach	Permission	Not ready	Consent state exposed in the service platform	High

Not which use cases are valuable, and not which are ready, but which valuable ones are ready, and precisely what stands between the rest and deployment. That overlay is where the roadmap comes from, and it is the subject of Paper 04.

What leaders should do now

Five actions, none of which require a data platform decision.

Test three workflows, not the estate

Take your three highest volume automation candidates and run the five tests against each. It takes under a week per workflow with the right people in the room, and the right people are not the data team alone. One person who owns the workflow operationally, one who knows the systems it touches, one who can speak to consent and retention.

Set a time budget for each one

Write down how current the data has to be for the decision, before anyone checks how current it actually is. Doing it in that order stops the available refresh rate from setting the standard.

Write the precedence rules

Start with the fields that appear in more than one system and are used in customer-facing decisions. Per field, with a divergence window.

Label outcomes for the workflows you are automating, not for the archive

Narrow and current beats broad and eventual, on a 14 month model lifecycle.

Sequence, do not abandon

A workflow that fails one test is a workflow with a date. Put the remediation on the plan and keep the use case in the portfolio.

Run the five tests against your own workflows

CXaiS assesses operational AI readiness across six dimensions, maps how work actually runs, and puts a commercial number on what is at stake before anything is built.

cxais.ai/contact

REFERENCE

About and sources

ABOUT CXAIS

CXaiS is an independent AI in CX consultancy. We build the commercial foundation, ROI framework and implementation plan that makes AI investment in customer experience perform. We are vendor independent and commercially accountable at every stage.

ABOUT THIS SERIES

CXaiS Institute publishes independent research on operational AI readiness for the enterprise. Insight Paper 01 sets out the six dimensions. Paper 02 sets out workflow cartography. This paper sets out the five tests of AI data readiness. Paper 04 puts a number on the value at stake. Read the series at institute.cxais.ai

DISCLAIMER AND SOURCES

Figures cited are drawn from published research as attributed. Observations described as ours are drawn from CXaiS client engagements and from the author's prior practitioner experience, and are stated as experience rather than research findings. This paper is published for information and does not constitute advice on any specific programme.

MIT Project NANDA, The GenAI Divide, 2025. RAND Corporation, The Root Causes of Failure for Artificial Intelligence Projects, 2024. S&P Global Market Intelligence, Voice of the Enterprise, AI, 2025. McKinsey, The State of AI, 2025. IBM Institute for Business Value, 2026.



cxais.ai · 2026