

PRODUCTION READINESS IN RETAIL

Is your data ready for AI?

An honest answer in three weeks.

Nearly every retail company has launched AI projects. Very few deliver in daily operations. This whitepaper shows why. And how you find out where things break down in your company and what to do about it.

PUBLISHER product GmbH · Unit [Elements]
VERSION July 2026 · Version 1.1 · free to share
READING TIME approx. 15 minutes

Why so much AI activity *achieves so little*

Few topics occupy retail as much as artificial intelligence. Chatbots in service, forecasts in buying, agents in replenishment: nearly every company has launched projects. Yet only a few report AI that works reliably in daily business and pays off. The numbers show both at once. Plenty of activity, little impact.

89.1 %

of retail companies name high data quality as a relevant success factor for AI. It is the top value in the survey.

[HDE AI STUDY 2025](#)

42 %

of companies abandon the majority of their AI initiatives before production. That is more than double the share of the previous year.

[S&P GLOBAL MARKET INTELLIGENCE, VOICE OF THE ENTERPRISE 2025](#)

> 40 %

of agentic AI projects will be cancelled by the end of 2027, Gartner forecasts. The reasons: unclear business value and missing controls.

[GARTNER · FORECAST JUNE 2025](#)

What is remarkable is *where* these projects fail. It is almost never the model. The [RAND Corporation](#) names unclear objectives and missing data as the most frequent causes. Not the technology. The [Cisco AI Readiness Index](#) regularly measures the data dimension as one of the weakest. And [Gartner warns](#) that a lack of AI-ready data puts projects at risk across the board. Three independent sources. One finding:

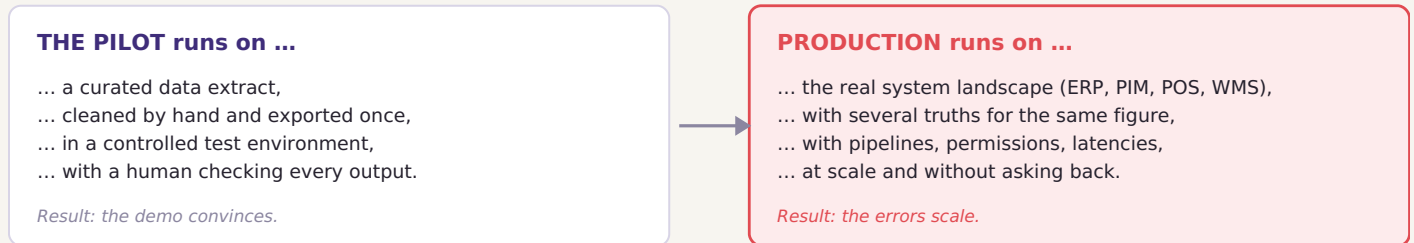
■ **The bottleneck lies in the data foundation.** It is about quality, accessibility, and ownership. Whoever starts the fifth pilot without knowing this foundation repeats the experiment. With the same outcome. The question is not whether your AI is good enough. The question is whether your data is good enough for the AI.

For decision makers this means: change the order. First measure whether the data is sufficient for the intended use case. Then invest. And invest where the measurement shows the bottleneck. What such a measurement looks like is the subject of the rest of this paper.

All figures come with source and year. The sources are linked at the end. Forecasts are marked as forecasts. We cite only robust surveys and leave out popular but disputed figures.

Why the pilot convinces *and* production disappoints

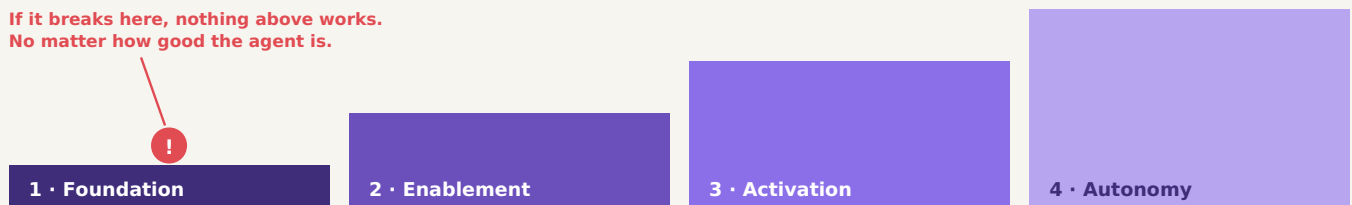
Why does the pilot convince while production fails? A pilot is an experiment under lab conditions. Someone selected the data, cleaned it, and exported it once. Production knows no such mercy. It runs on the real system landscape. With everything that has accumulated there over the years:



A human copes with faulty master data. He spots the error, asks a colleague, or quietly corrects it. An AI system does not. It processes what is there. In seconds, at scale, without asking. **AI does not solve data problems. It scales them.** One example: if an item is missing its lead time, the agent guesses on every order. Just automated, a thousand times over.

The weakest layer limits the whole

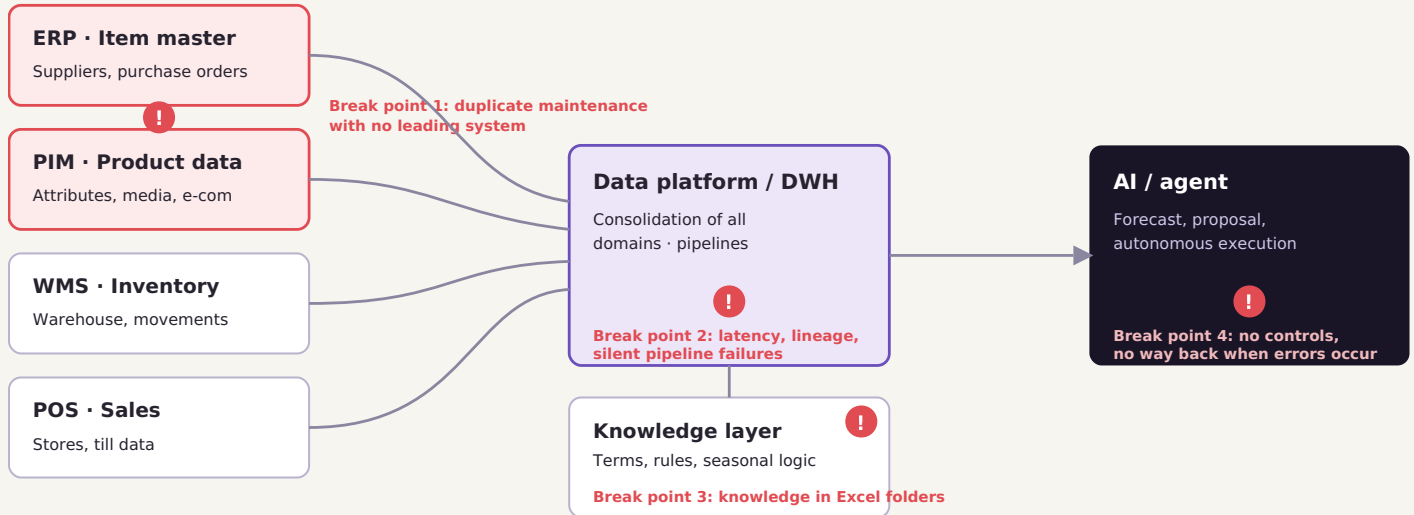
Every AI use case rests on four stages. The **Foundation**: clean data with clear ownership. The **Enablement**: data becomes usable through pipelines and interfaces. The **Activation**: company knowledge becomes retrievable for AI. The **Autonomy**: agents carry out tasks on their own. You can skip stages. But not without consequences. The weakest required layer limits how reliable the whole can become. No agent tuning lifts this limit. Only closing the gap does.



This one layer is what we call the **bottleneck**. Finding it is the purpose of the production readiness audit. Because it shows where the next investment belongs. Almost never into more AI. Almost always into one specific, nameable data layer.

Where the bottleneck *arises in your data flows*

In retail, an AI decision never comes from a single system. An order proposal is the end of a chain. Master data comes from ERP and PIM. Stock levels from the warehouse. Sales from the till. Terms from the supplier. A data platform brings it all together, and the AI works with the result. Every link in this chain can be the bottleneck. And the typical break points look similar from company to company:



The four break points match the four stages. Duplicate maintenance and missing ownership are Foundation problems. Pipeline failures and latency belong to Enablement. Unmaintained knowledge is an Activation problem. Missing controls an Autonomy problem. **An audit checks exactly this chain. But only the links your use case actually needs.** That is what separates a targeted review from a company-wide data project.

What separates an honest assessment *from a questionnaire*

The market is full of free readiness checks based on questionnaires. Their problem: self-assessments come out too positive. Every department rates its own work. A score without evidence is an opinion with a decimal place. An assessment you can base an investment decision on follows three principles:

PRINCIPLE 1

Evidence or capped

Every criterion needs a verifiable artifact. A data quality report, interface documentation, an operations log. Without evidence there is no good grade. No matter what is reported.

PRINCIPLE 2

Per use case, not blanket

Not: how mature is the company? But: which data layers does this use case need? And are they good enough for it? What is measured is a matrix of stages and data domains (see below).

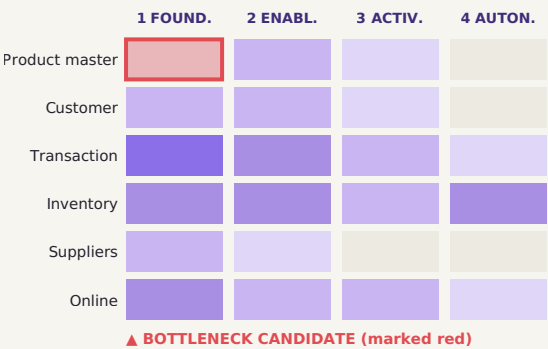
PRINCIPLE 3

Reproducible

Criteria, scoring anchors, and formulas are openly documented. Two assessors must reach the same result with the same evidence. Otherwise it is not a standard but gut feeling.

What is measured is a heatmap, not a single number

Every combination of data domain and stage is scored. The domains in retail: product master data, customer, transaction, inventory, suppliers, online behavior. Each cell receives a maturity score from 0 to 100. A single overall number would deceive, because maturity is unevenly distributed in practice. The heatmap shows where your data is strong. And where exactly it jams:



The heatmap is read from the use case. A reordering agent needs product master and inventory. It does not need the online domain. Only the cells the use case requires are scored and closed.

The scale has five maturity bands: Ad hoc, Initial, Established, Managed, Optimized. Every result is given as a band with confidence. Never as a pseudo-precise single number.

One assessment dimension is often forgotten: can AI agents even reach the data? Agents work through interfaces, not Excel exports. So we test this technically. Are there APIs? Are they documented? Are the data models machine-readable? This check reads the systems themselves. Interviews are not enough for that.

Illustrative heatmap. The full scoring logic (criteria catalogs, maturity bands 0-100, target derivation, formulas) is documented in the open measurement standard: prodct.group/foundation-ascent/methode

What the audit puts *in your hands*

An audit does not end with a slide deck full of recommendation prose. It ends with a finding and four results. Each is usable on its own. Together they are the basis for your next AI investment:

R1

The reliability score

How reliable can your use case become today? An evidenced number from 0 to 100, given as a band with confidence. Not an opinion. It shows whether the targeted level of automation is realistic today.

R2

The bottleneck

Which data layer limits the result? With evidence from your systems and the price of the gap. The bottleneck answers the question: where to invest first?

R3

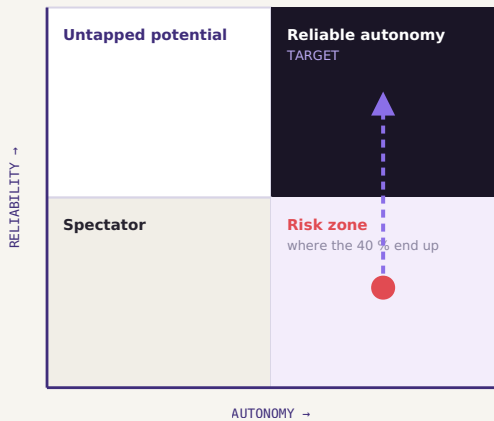
The action plan

What needs to happen first? What does it cost, what does it deliver? The plan is prioritized and described so that you can execute it without us.

R4

The positioning

Where do you stand in the target picture of reliability and autonomy? Plus an honest recommendation for the next step. Even if it reads “no agent yet”.



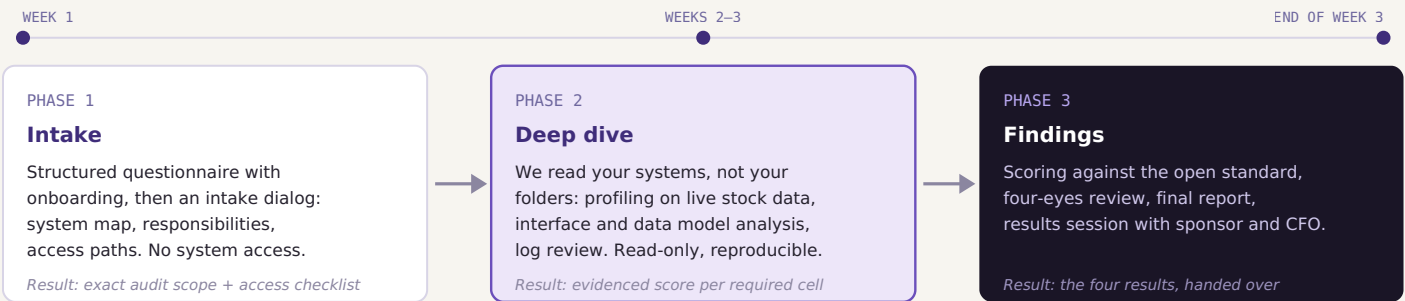
The way forward is up, not to the right.

Role	What the audit gives you
COO / Business unit	The answer why the process does not yet run reliably. And a plan with an expected effect per measure. Instead of yet another pilot.
CDO / CIO	Evidence of which data layer is fine and which one slows you down. Measurable instead of gut feeling. Plus a clear investment sequence to support budgeting.
CFO / Management board	Fixed price, fixed end date, open-ended finding. And the certainty of not putting more money into a use case that another layer is limiting.

■ The red dot in the quadrant is the most common finding: high autonomy ambition on a weak foundation. That is the direct route into the cancellation statistics. Nobody needs autonomy for its own sake. What counts are processes that run reliably. With as much autonomy as the foundation allows. **Reliability before autonomy.**

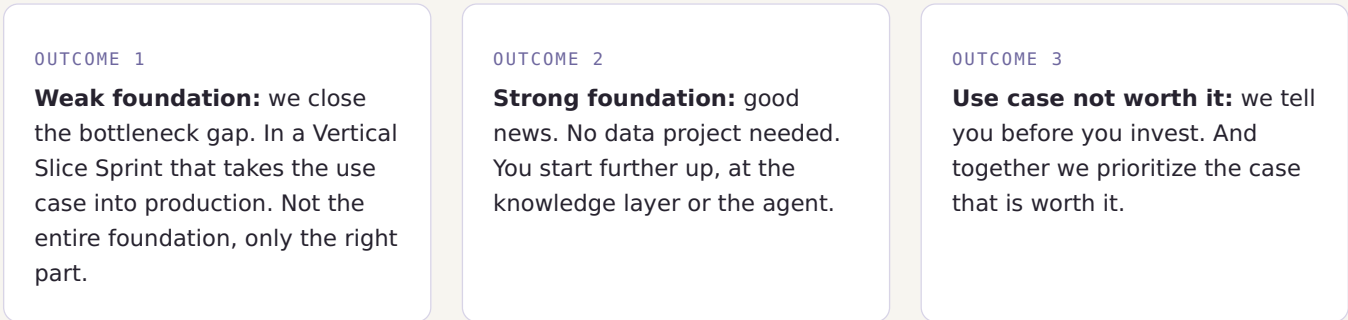
How the audit runs in three weeks

An audit does not have to paralyze operations. The process therefore has two parts. First we take stock in a structured way, without any system access. Then we examine, read-only, the systems your use case needs. Your effort stays small and the depth of the review stays high:



Your effort stays manageable. One half-day workshop. One contact person with two to three hours per week. Read-only access per checklist. No preparation, no cleanup beforehand. The audit assesses the current state, that is exactly what it is for. If access is missing, the clock pauses transparently.

The audit is open-ended and has three possible outcomes



■ All results belong to you and can be executed without us. An uncomfortable answer after a three-week audit is cheaper than twelve months of pilot without production.

A look into a finished report, *the ACME Inc. example*

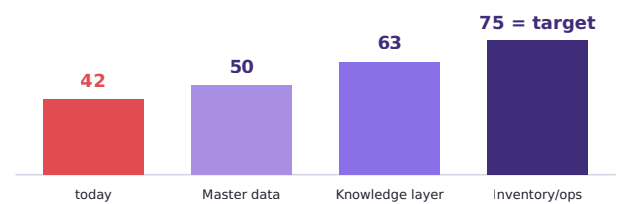
How concrete a report becomes is best shown by example. ACME Inc. is fictitious but fully worked through. A multichannel retailer with around €480 million in revenue wants to automate reordering. The agent is meant to trigger order proposals on its own. The planner, the person who decides on replenishment orders, only steps in for exceptions. The audit does not find the bottleneck in the agent. It finds it in the item master. It is maintained in ERP and PIM in parallel, with no leading system. The consequences are measurable: duplicates, missing lead times, missing minimum order quantities. Exactly the fields an ordering decision needs.

THE FINDING IN NUMBERS (FICTITIOUS)

Required cell	Target	Actual	Gap
Product master · Stage 1	75	42	33 ◀ BOTTLENECK
Inventory · Stage 1	75	63	12
Inventory · Stage 2	75	54	21
Inventory · Stage 3	75	50	25
Inventory · Stage 4	75	55	20

Reliability score today: **band 37-47 out of 100**. Too low for the target level “autonomous with exception approval” (target 75). No agent tuning lifts this score. Only closing the master data gap does.

THE EFFECT OF THE ACTION PLAN (FICTITIOUS)



The action plan lifts the score step by step. Every measure has an effort and an expected effect. In the end the data base is sufficient for the target.

■ The full example report has eleven pages. It contains a management summary, system landscape, data model with measured field completion, bottleneck evidence, action plan, and positioning. A dossier makes every number traceable. Free to download: prodct.group/foundation-ascent/produktionsreife-audit

ACME Inc. is fictitious. Worked example, fictitious data: the numbers demonstrate the methodology. Real audits are confidential. Only anonymized, aggregated values flow into the benchmark.

CONCLUSION

Why your AI initiative really fails *can be found out*

Three sentences sum up the situation. Most AI initiatives in retail do not fail because of the model, but because of the data foundation. Which layer limits a use case can be measured in three weeks. With evidence and against an open standard. And whoever knows the bottleneck invests differently: precisely into a nameable layer instead of broadly into more AI.

So the question is not whether you will work with AI. The question is whether your data is ready for it. Both start in the same place: with an honest measurement.

NEXT STEP

Request a production readiness audit

Three weeks, fixed price, open-ended. An answer within 24 hours. From a person, not a funnel.

PRODCT.GROUP/FOUNDATION-ASCENT

THE METHOD IN DETAIL

Open measurement standard

Criteria, maturity bands, and formulas are publicly documented. Traceable instead of a black box.

PRODCT.GROUP/FOUNDATION-ASCENT/METHODE

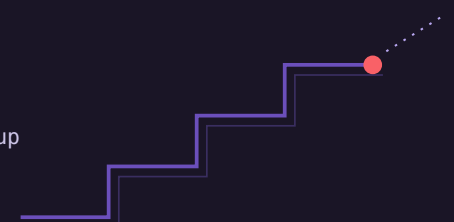
THE PROOF

See the example report

The full ACME report as a PDF. This is what the result you hold after three weeks looks like.

FREE TO DOWNLOAD ON THE AUDIT PAGE

CONTACT Jan Fischer · jan.fischer@prodct.group · WEB prodct.group
Foundation Ascent® is a method by prodct GmbH · Unit [Elements]



The sources of this paper, *linked and put in context*

Many figures about failing AI projects are in circulation. They differ considerably in method. We cite only surveys with a traceable origin and methodology. Forecasts are marked as forecasts. We leave out popular but disputed figures.

Source	Finding & context
HDE · AI study in retail 2025 (with Safaric Consulting)	89.1 % of retail companies name high data quality as a relevant or very relevant success factor for AI, the highest value in the survey. The most retail-specific source in this paper.
S&P Global Market Intelligence · Voice of the Enterprise: AI & Machine Learning, Use Cases 2025	The share of companies abandoning the majority of their AI initiatives before production rose from 17 % to 42 % (n = 1,006, North America and Europe). The most solid cancellation figure on the market.
Gartner · Forecast on agentic AI projects, June 2025	Over 40 % of agentic AI projects will be cancelled by the end of 2027. The reasons: rising costs, unclear business value, missing risk controls. A <i>forecast, not a measured result</i> .
Gartner · “Lack of AI-ready data puts AI projects at risk”, February 2025	63 % of data management leaders do not have suitable data practices for AI or are unsure whether they have them. A lack of AI-ready data puts projects at risk across the board.
RAND Corporation · The Root Causes of Failure for AI Projects, 2024	Interviews with 65 experienced data scientists and engineers. Most frequent causes: misunderstood objectives and missing data. Not the technology. A qualitative study.
Cisco · AI Readiness Index 2024	Survey of 7,985 executives in 30 markets. Only 13 % of companies are fully ready. The data dimension is regularly among the weakest pillars.
Informatica · CDO Insights 2025	Survey of 600 data leaders. Data issues are the most frequent hurdle on the way from pilot to production. A vendor-affiliated source. The figures match the independent surveys.

■ **What we deliberately do not cite:** the widely quoted “95 % failure rate” (MIT/NANDA 2025). It is narrowly defined, methodologically disputed, and usually misquoted. Whoever argues with numbers must also be able to let go of bad ones. That holds for studies as it does for master data.

Links checked: July 2026. © prodct GmbH. This whitepaper may be shared freely in unchanged form.