

PROPOSAL

Where AI should keep, help, or replace human work in *retail*.

Most retailers now have a list of AI ideas and no reliable way to tell which ones will pay. Two things fix that: a method for sorting use cases by how much human judgement they should keep, and an honest read of how ready your people are to work alongside AI. This is both.

The cost of getting this wrong is not the failed pilot. It is the two years spent on the wrong ones while a competitor quietly automates the right ones. In earlier work the same discipline cut reporting time by 90% across 56 countries and took a consolidation from three hours to three minutes — not because the tools were clever, but because the right tasks were chosen first.

15+ years

in FMCG and retail — value engineering, commercial excellence, supply chain

Data & AI consulting

through Data Lab: AI strategy, analytics, working tools across Europe

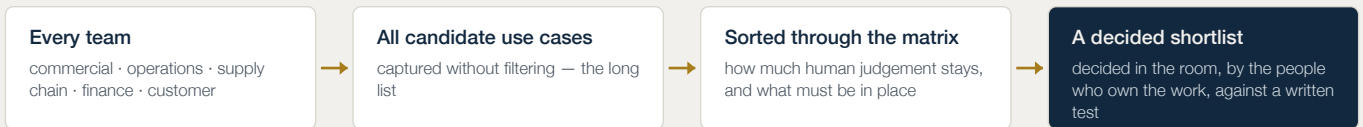
Executive teaching

AI & Strategy for Executives at SSE Riga and Hanken & SSE; 3,000+ people trained

WHY ME FOR RETAIL

Before Data Lab I spent fifteen years inside Carlsberg Group, mostly on the commercial side of retail. In Norway I ran value engineering on price, promotion and assortment at Ringnes, the country's largest soft-drinks business — the three levers where FMCG margin is won or lost at the shelf. Across France, Croatia, Italy, the Baltics and Sweden I built value engineering and sales-force capability in markets with very different conditions. At head office I was the supply-chain business partner and, later, internal consultant for commercial excellence and insight from data. I know how a retailer's numbers move, because I moved them.

WHAT WE WANT TO DO



METHOD · 1

Keep human · Augment · Automate

Every candidate use case is placed on one of three tiers. The question is never 'can AI do this?' — it is 'how much human judgement should stay, and what has to be in place for that to work safely?'

□ more human judgement kept

more infrastructure required □

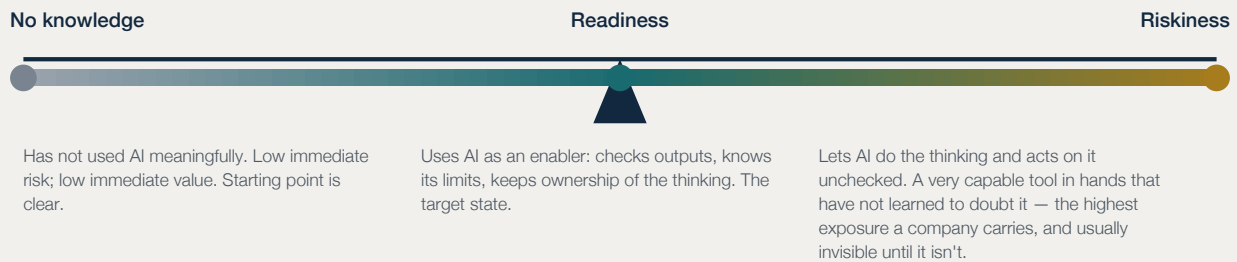
KEEP HUMAN	01 AI not capable yet The task needs judgement, context or trust the tools cannot yet provide reliably. Revisit in twelve months.	02 Too high risk Regulatory, safety, reputational or ethical exposure outweighs any gain. A human owns the outcome.	03 Not worth it AI could do it, but the value is too small to justify the change, the risk, or the attention.
AUGMENT	01 Human in the loop AI drafts, suggests, or scores; a person decides every time. Enablement: good tools, prompts and training — the process itself barely changes.	02 Human on the loop AI acts within set bounds; a person supervises, samples and intervenes. This needs process redesign: new roles, escalation paths, and clear thresholds.	03 Human over the loop AI runs the work; a person owns outcomes and direction, not individual cases. This needs governance, monitoring, audit trails, and a defined way to pull the plug.
AUTOMATE	01 Simple tasks Repetitive, low-stakes, well-defined. Minimal change management; quick wins that build confidence.	02 High-effort, lower-value Important but tedious work that consumes skilled time. Automate to free people for judgement work.	03 High-impact tasks Decisions that move revenue or cost. Requires the most: clean data, integrated systems, risk management, monitoring, and a tested rollback.
WHAT EACH TIER NEEDS	A named owner, a review date, and a written reason — so 'keep human' is a decision, not a default.	Role clarity, decision rights, and the training to match — augmentation fails quietly when nobody knows who is accountable.	Data quality, system integration, monitoring and ownership in place before go-live — not promised for afterwards.

For example: promotional uplift forecasting usually lands in Augment, human in the loop. Shelf-gap detection from photos lands in Automate, simple task. Markdown pricing on fresh goods stays Keep human, too high risk — for now.

METHOD · 2

AI Readiness & Risk — the Balance Scale

Use cases only work if the people running them are ready. The Balance Scale gives each person — and the organisation — a reading on one line: from no knowledge, through readiness, to riskiness. It scores a pattern of use, not a person; the same individual can sit in different places for different kinds of work.



HOW IT IS MEASURED

01 Self-assessment Where people think they are. Quick, and revealing when compared to the rest.	02 Quick cases Short, realistic decisions with AI in the mix. Do they check, question, and verify?	03 Situational tasks What would you do if the model were confident and wrong? Tests judgement under ambiguity.	04 Behavioural tasks Observed use of real tools on real work. The gap between stated and actual habits.	05 Exercises Guided work that both measures and teaches — so the assessment is itself a first step.
---	--	--	---	---

The point is not to grade people. It is to know who is ready to be trusted with which tier of use case — and who needs support before they are. Most organisations are surprised by where their readiness actually sits.

APPROACH

Discovery, assessment, roadmap.

01 Discovery

We sit with the teams and capture every candidate use case, then place each one on the matrix.

02 Assessment

We read readiness and risk across the people who would run them.

03 Roadmap

You get a ranked shortlist with the reasoning written down — what to do first, and what to leave alone.

Typically a few weeks and a handful of working sessions with the teams who own the work. Scope follows the size of the organisation; we agree it together.

Rihards Garančs

Founder & strategy consultant, Data Lab SIA · Strategic Advisor, SSE Riga Executive Education

rihards.garancs@gmail.com

rihardsgarancs.com/company · Riga, Latvia