

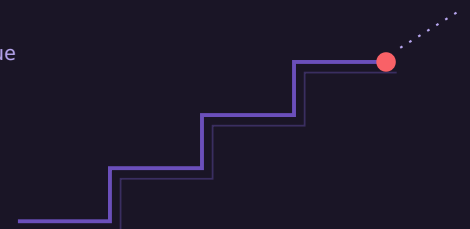
PRODUCTION READINESS AUDIT · FINDINGS REPORT

ACME Inc.

Agentic replenishment.

How reliable can automated replenishment at ACME become today? What is holding it back? And what does it take to run it in daily operations? This report answers these three questions with evidence.

CLIENT	ACME Inc. (fictitious) · multichannel retail · ~€480M revenue
PERIOD	June 29 – July 17, 2026 (3 weeks incl. intake)
STANDARD	Measurement Standard V1.0 (public) · Assessment Form V1
ASSESSOR	[Lead Assessor] · REVIEW [Reviewer], July 12, 2026



The result *on one page*

Your use case: automated replenishment for stores and the online shop. An AI agent is meant to trigger orders on its own. Your planner, the person who decides on replenishment orders today, only steps in for exceptions. **Your target metric:** the out-of-stock rate. It measures how often an item is not available. Today 6.8 % (source: business unit, workshop July 1, 2026).

R1 · RELIABILITY SCORE

37-47 /100

Band "Established" · confidence: high

R2 · BOTTLENECK

Product master data · Stage 1

Actual 42 of required 75 · gap 33

R3 · ACTION PLAN

5 actions

First: designate the leading system for product master data (effort M)

R4 · ASSESSMENT

Heading toward the risk zone

Recommendation: close the bottleneck in a Vertical Slice Sprint (10-14 weeks)

■ **The essentials in three sentences:** Your replenishment agent can reach a reliability score between 37 and 47 out of 100 today. Autonomous ordering with exception approval requires 75. The cause is not the AI but the product master data. It is maintained in two systems in parallel, with neither in the lead. The very fields an order depends on are therefore often missing (lead time, minimum order quantity). Closing this gap lifts the score to 50. The full plan takes it to 75. Then your data foundation is sufficient for the goal.

All scores collected on an evidence basis (rule: evidence or capped), four-eyes review on July 12, 2026. Every figure in this report can be traced through the cell dossier (Appendix A). This report is an anonymized example with fictitious data.

Scope *and benchmark*

An audit never measures the whole company. It measures whether your data is good enough for one specific use case. That requires two things: a target (how reliable does it have to be?) and the list of data domains that matter.

REQUIREMENTS PROFILE (WORKSHOP JULY 1)

Automation level: **autonomous with human-in-the-loop**, meaning the AI acts and a human approves exceptions
Cost of errors: **high**. Wrong orders go to real suppliers and tie up capital
Regulatory sensitivity: no

TARGET DERIVATION (MEASUREMENT STANDARD §4.3)

Base Autonomous+HITL 65

Error costs
high +10

TARGET/cell
75

The required cells. A cell is the combination of a data domain and one of the four stages. Your use case needs these five:

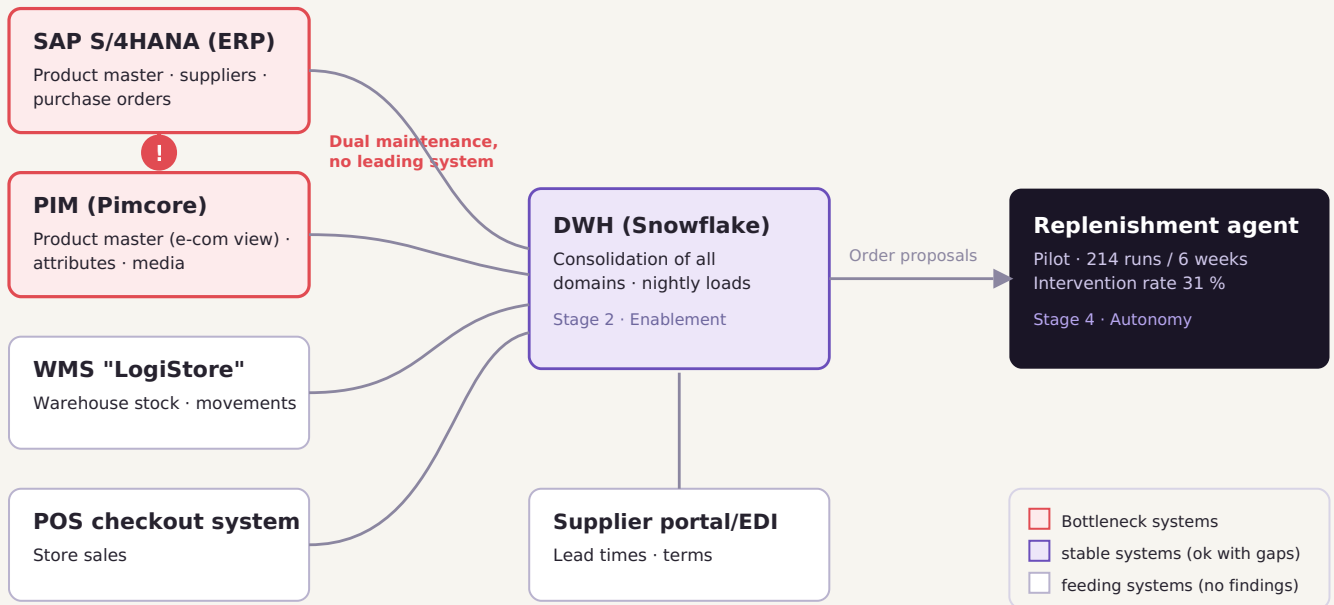
Cell (domain • stage)	Why required
Product master data • Stage 1	Every ordering decision depends on item attributes: lead time, minimum order quantity, pack size, supplier assignment.
Inventory/supply chain • Stage 1	Inventory accuracy is the second decision basis of every replenishment order.
Inventory/supply chain • Stage 2	The agent needs inventory as a dependable, up-to-date data flow. A manual export is not enough.
Inventory/supply chain • Stage 3	Retrievable, verified knowledge: supplier terms, assortment rules, seasonal logic.
Inventory/supply chain • Stage 4	The existing pilot agent: success rate, interventions, control in (test) operation.

DATA COLLECTION (JULY 2–8)

23 artifacts reviewed. An artifact is a verifiable piece of evidence, such as a report or system documentation. Plus our own profiling, an automated data analysis run directly in the systems: the complete product master data (84,312 active items in SAP and PIM), inventory as a sample (12 stores and the central warehouse). In addition, the pilot agent's operating logs (6 weeks, 214 runs) and 4 interviews. Scored against the open Measurement Standard V1.0: prodct.group/foundation-ascent/methode.

System landscape *and bottlenecks*

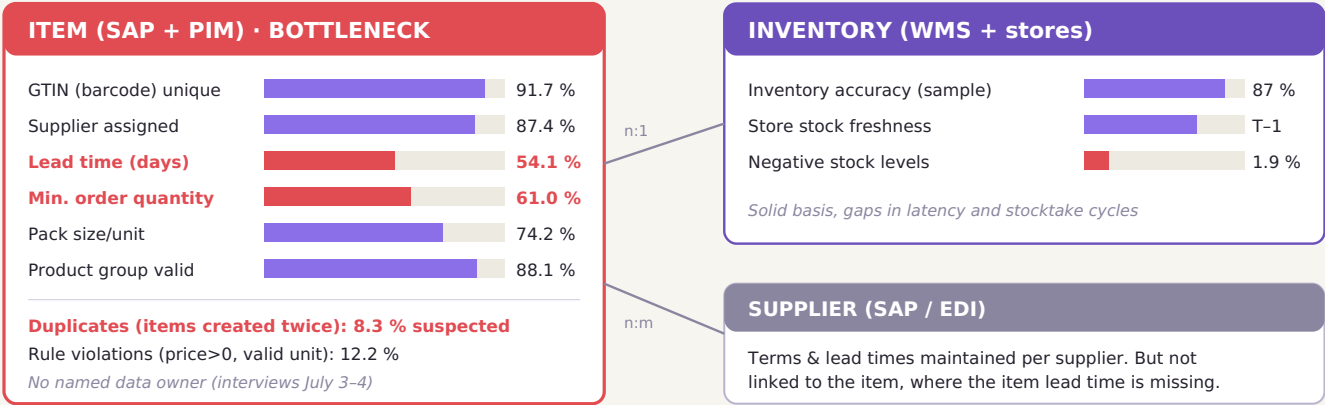
Seven systems feed data into replenishment. SAP is your ERP, the central system for merchandise management and purchasing. The PIM maintains product data for the online shop. Both maintain the product master data in parallel, **with neither system in the lead**. That is exactly where the bottleneck arises (marked in red).



■ **Finding:** The data warehouse (DWH), the central database that brings everything together, works reliably. The problem arises upstream. Two systems claim the truth about the item. Their contradictions travel unfiltered into the agent's order proposals.

The state *of your data*

The three most important data domains of your use case, with the fields that matter for an order. The percentages show for how many of your 84,312 active items the field is maintained. Measured directly in the systems on July 7, 2026.

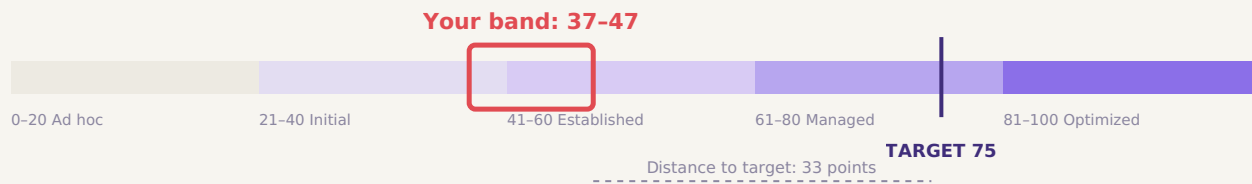


Why exactly this slows the agent down

A replenishment decision depends on two item fields above all: lead time and minimum order quantity. Exactly these are the weakest maintained (54 % / 61 %). For almost every second item the agent has to guess or a human has to step in. That explains the measured intervention rate of 31 % in the pilot (logs, Appendix A).

Reliability score *in the lower mid-range*

The score is the lowest value of your five required cells. Not the average. Because the weakest point determines how reliable the whole becomes. We state the score as a band with confidence. Confidence says how solid the evidence behind it is.



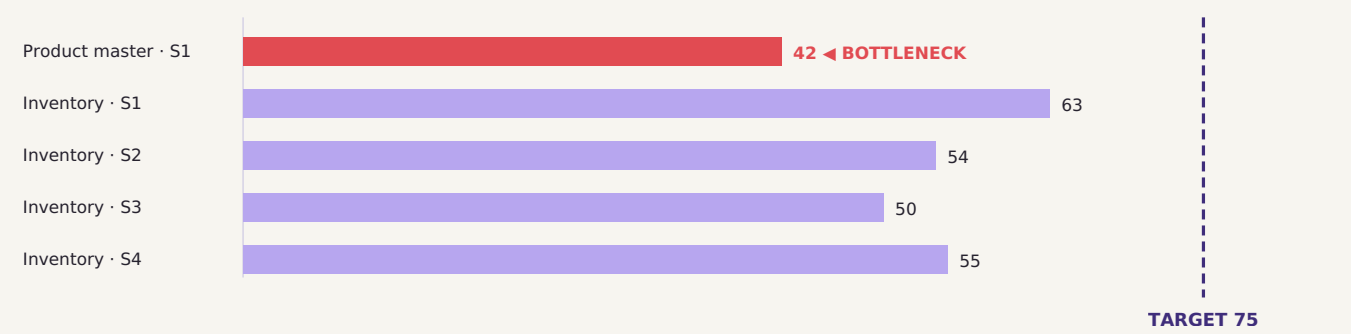
Required cell	Target	Actual	Band	Confidence	Gap
Product master data · Stage 1	75	42	Established	high	33 ◀ BOTTLENECK
Inventory/supply chain · Stage 1	75	63	Managed	high	12
Inventory/supply chain · Stage 2	75	54	Established	high	21
Inventory/supply chain · Stage 3	75	50	Established	medium	25
Inventory/supply chain · Stage 4	75	55	Established	high	20

WHAT THIS MEANS IN OPERATION

An agent on this basis makes wrong or incomplete order proposals for many items. A human has to supervise it constantly. Your pilot confirms this: someone had to intervene in 31 % of proposals. For autonomous ordering with exception approval (target 75) that is not enough. **Important: This score cannot be improved on the agent itself. Only at the bottleneck.**

The identified *bottleneck*

The cell product master data at Stage 1 reaches 42 of the required 75 points. All other required cells are clearly higher:



How we substantiate this. All evidence is in Appendix A:

Finding	Evidence (artifact, date)
The mandatory fields for replenishment are only 71 % filled. Critical: lead time 54.1 %, minimum order quantity 61.0 %	Profiling P-01, Jul 7, n=84,312
Suspected duplicates (items created twice) at 8.3 % . The cause is dual maintenance in SAP and PIM	Profiling P-02, Jul 7
Rule violations 12.2 % (price > 0, valid unit/product group)	Rule set R-01 + profiling P-03
Nobody is officially responsible for the product master data. Purchasing and e-commerce correct in parallel	Interviews I-02/I-03, Jul 3–4
In the pilot agent a human had to intervene in 31 % of runs. 72 % of these interventions were related to master data	Agent logs L-01, 6 weeks, 214 runs

WHAT THE BOTTLENECK COSTS TODAY

Your business unit named it in the workshop: order proposals have to be corrected by hand. Items without a lead time run out of stock. Extra stock is ordered to be safe, which ties up capital. We deliberately do not put a number on this. The real value shows in operation, in the sprint's 90-day measurement.

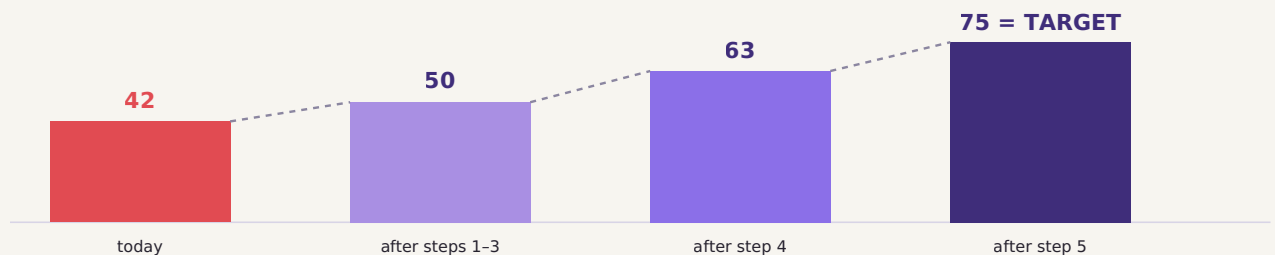
WHAT CLOSING IT COSTS

Gap	Effort	Person-days
Product master · S1 (33)	M	35-55
Inventory · S2/S3 (21/25)	M	25-40
Inventory · S1/S4 (12/20)	S	10-20

The action plan *in five steps*

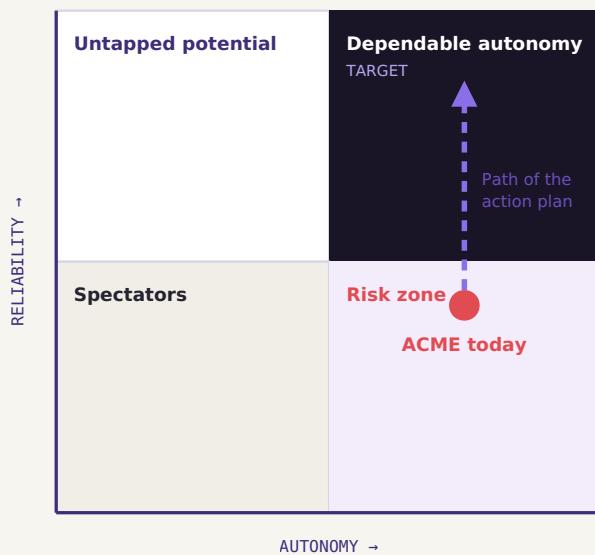
The order follows the benefit: what lifts the score most comes first. Each action is described so that ACME can implement it without prodct. The column Effect after shows the new reliability score once the action takes effect.

#	Action	Closes	Effort	Effect after
1	Decide which system leads the product master data. Our recommendation: SAP leads, the PIM follows. Define reconciliation rules and clean up the duplicated items (suspect list is attached)	Product master S1	M	Bottleneck closed → score jumps to 50 ; then Inventory-S3 becomes the limit
2	Fill the mandatory fields for replenishment. Start with lead time and minimum order quantity for the top-selling items (62 % of order volume). The values come from the supplier portal, not from manual entry	Product master S1	M	
3	Assign ownership. One named person is responsible for the product master data. New and changed items pass automatic rule checks	Product master S1	S	
4	Organize the knowledge. Supplier terms, assortment and seasonal rules move from Excel folders into a verified, versioned source the agent reads from	Inventory S3 + S2	M	Score jumps to 63 ; then Inventory-S1 becomes the limit
5	Make inventory more accurate. Cycle counts for the most important items, clear rules for negative stock, store inventory updated several times a day instead of once at night	Inventory S1 + S4	S-M	Score reaches 75 = target. The data foundation is sufficient for the goal



■ **What we deliberately do NOT recommend:** further tuning of the agent, a new PIM project, or a company-wide data cleanup. Only what limits this use case gets closed. The smallest foundation that is sufficient.

Assessment *and recommendation*



Your position: The autonomy ambition is high, the measured reliability is 37 to 47. That places ACME at the edge of the risk zone. This is where the projects land that later get cancelled: high autonomy on a weak foundation. The good news: the way up is short and clearly described.

OUR RECOMMENDATION · OPTION 1

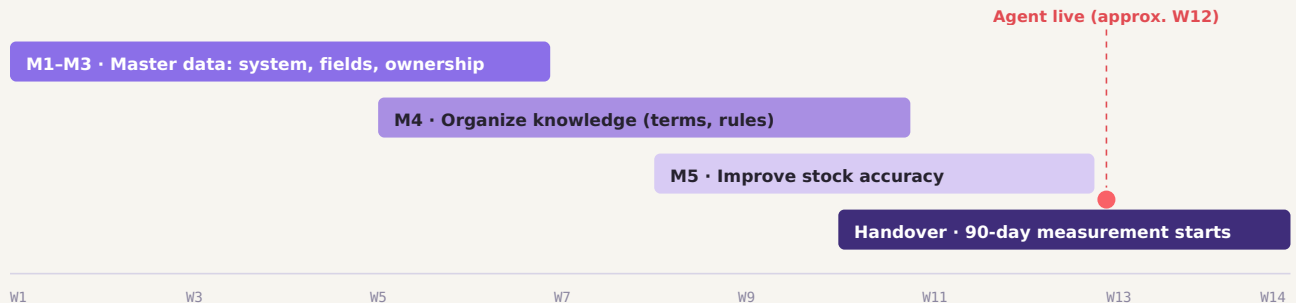
Close the bottleneck in a Vertical Slice Sprint. The sprint takes a single use case across all four stages into real operation. Actions 1 to 3 are the core, 4 and 5 run in the same pass. At the end the agent works in production, not in another test. Duration: **10 to 14 weeks**. Investment: **[range, proposal to follow]**. Done means: production operation plus measurement against your target metric (out-of-stock rate 6.8 %) after 90 days of operation. That is written into the contract.

YOUR NEXT STEPS · EVEN WITHOUT US

- 1 · Decide which system leads the product master data. This is a policy decision, not a project.
- 2 · Have purchasing and e-commerce review the attached list of suspected duplicates.
- 3 · Take lead time and minimum order quantity for the most important items from the supplier portal.

The road *to implementation*

A finding is only worth something once it is implemented. The roadmap below shows how the Vertical Slice Sprint closes the bottleneck and takes the agent into production. It works with us or with your team. The sequence stays the same.



Corridor of 10 to 14 weeks. Done means: the agent works in day-to-day business, and the 90-day measurement of your target metric (out-of-stock rate, today 6.8 %) is running. That is fixed in the contract.

WHAT ACME TAKES ON

- 1 · Policy decision: which system leads the product master data (management and IT).
- 2 · Review the duplicate candidate list (purchasing and e-commerce, list enclosed).
- 3 · Appoint a data owner for the product master data and give them time.
- 4 · One contact in the business, two to three hours per week.

WHERE PRODCOT SUPPORTS

- 1 · Build synchronization rules and automated rule checks (M1, M3).
- 2 · Pull lead times and minimum order quantities from the supplier portal, automated (M2).
- 3 · Set up the knowledge layer as a verified, versioned source (M4).
- 4 · Take the agent into production with logging, limits and exception approval (M5).
- 5 · Set up and evaluate the 90-day value measurement.

THE FIRST 14 DAYS · READY TO START, EVEN WITHOUT US

Week 1: bring about the decision on the leading system and appoint the data owner. Both are decisions, not projects. · Week 2: have purchasing and e-commerce review the duplicate list, and pull the lead times for your A items from the supplier portal. That starts measures 1 to 3 before the sprint begins.

Derivation

Each cell is scored against fixed criteria with 0 to 4 points. The points produce the cell value. The complete dossier is in your project folder. Here is the extract for the bottleneck cell, with the remaining cells in short form.

Product master data · Stage 1 · cell value 42 (= 100 × 10 ÷ 24) · confidence high

Criterion	Points	Evid.	Artifact (date)	Rationale
1.1 Completeness	2	Y	P-01 profiling Jul 7	Mandatory fields 71 % filled, critical fields well below. Maintenance is selective, not to a standard.
1.2 Uniqueness	1	Y	P-02 duplicates Jul 7	8.3 % suspected duplicates. Cause is dual maintenance in SAP and PIM, without control.
1.3 Consistency	2	Y	R-01 + P-03	A rule set exists (purchasing) but is not checked systemically; 12.2 % violations.
1.4 Timeliness	2	Y	P-04 timestamps	New items current; existing items without change-date maintenance in PIM.
1.5 Ownership	1	Y	I-02/I-03 interviews	No named owner; purchasing and e-com correct in parallel, person-dependent.
1.6 Compliance	2	Y	D-05 deletion concept	Concept exists, implementation not evidenced.

Remaining cells (short form)

Cell	Value	Confidence	Key finding
Inventory · S1	63	high	Inventory accuracy 87 % (sample of 12 stores + central warehouse); negative stock 1.9 %; stocktake cycles irregular.
Inventory · S2	54	high	Nightly loads stable (error rate 1.2 %, 90 days of logs); lineage only partly documented; store latency T-1.
Inventory · S3	50	medium	Terms and assortment rules live in Excel folders, in several versions. Knowledge test with 20 domain questions to the AI: 12 correct, 5 without a source, 3 wrong.
Inventory · S4	55	high	Pilot instantiated: success rate 69 %, interventions 31 %; monitoring in place, rollback defined, governance approvals incomplete.

Confidence medium for Inventory-S3: the index of knowledge sources is incomplete. This has no effect on the bottleneck finding, the bottleneck cell has high confidence. Re-assessment planned in the sprint.

Methodology, limits *and terms*

METHODOLOGY

The basis is the open Measurement Standard V1.0. Each cell has 4 to 6 criteria. Each criterion receives 0 to 4 points, with a fixed description per score. The rule is: evidence or capped. Without verifiable proof the score stays low. The cell value is $100 \times \text{points achieved} \div \text{points possible}$. The reliability score is the lowest value of all required cells. The target follows from the automation level and the cost of errors. Every finding passes a four-eyes review. The standard is public: [prodct.group/foundation-ascent/methode](#).

LIMITS · HONESTLY

The score is a justified upper bound. It says: with this data foundation, the use case will not become more reliable than this band. It does not say the band is reached automatically. Actual reliability only shows in operation, measured over 90 days in the sprint. Confidence: high means every criterion is backed by original evidence. Medium means at least 60 %. Low is below that.

TERMS

Bottleneck · the data layer that limits how reliable your use case can become (technical term: reliability ceiling). **Cell** · the combination of a data domain and a stage, scored from 0 to 100. **Gap** · the distance between target and actual for a cell (technical term: foundation debt). **Confidence** · how solid the evidence behind a score is. **Profiling** · automated data analysis directly in the systems. **Planner** · the person who decides on replenishment orders. **Vertical Slice Sprint** · one use case across all four stages into production operation. **The four stages** · Foundation, Enablement, Activation, Autonomy.

NEXT STEP

In the results meeting we discuss the sprint proposal based on these findings.

All results in this report belong to ACME Inc. and can be implemented without prodct. Contact: Jan Fischer · jan.fischer@prodct.group · [prodct.group](#)

Example report with fictitious data (ACME Inc. does not exist). Figures consistent with the reference calculation example of Assessment Form V1. Design colors (Amethyst/Unit Elements) preliminary, final from the prodct design system.