



Designing AI-native experiences: The rise of human-agent interaction

Cross-referencing with Q4 benchmarks...



Worth flag

A 12% margin drop in the target segment. Want to get closer?

Dismiss

Action Required

I've drafted a response to the project delay. How would you like to proceed?



Foreword: Objective-first, not workflow-first

The next transformational shift in technology is here. The last time we witnessed an evolution of this speed was the arrival of personal computers—decades ago. This is artificial intelligence (AI), but not as we've known it.

In cases of traditional automation and personalization, it has always been humans who have defined the steps. In contrast, the very nature of agentic AI means that we may decide the end goal, but certainly not the journey. There is a sense that in order to maintain control, we're required to relinquish it.

This new operating model calls for a mindset shift, from workflow-first to objective-first. As AI agents create a new layer between human and system, and control becomes compromised, the focus for organizations shifts from merely designing workflows to designing boundaries and constraints. From learning how to interact with machines, to designing systems that can comprehend and operationalize human intent.

This next wave of technology is democratizing opportunities for people and society, and it's both liberating and overwhelming. Humility and empathy are vital; AI delivery must make sense for the business, the domain and the risk profile. If companies can't maintain objectivity, if they're not able to realize actual lasting value beyond output and efficiency gains, they'll jeopardize trust and ultimately adoption. Those who can step back and take a systemic view while staying anchored by their values will be the ones who are able to transform with confidence.

This report will walk you through our own approach to how we enable organizations to be able to do this, and how people and systems should interact to ensure the safeguarding of human integrity, encompassing AI-first delivery and depth of industry expertise. It informs our vision for a collective set of principles that sees that humans always stay in control, no matter how advanced AI gets.

Michael Schreibmann, CEO & Co-founder, Star



Introduction:

Post-chat: the new AI interaction paradigm

The first, chat-based AI interaction model was predictable, novel and simple; you'd type in a prompt and wait for a response. This model has been enough for a slew of companies to radically overhaul processes and teams, in some cases to the point of obsolescence.

In recent months, agentic systems have catapulted us into a post-chatbot era of AI. It marks a departure from the linear, conversational formula we've come to associate with large language models (LLMs), and an ascension into autonomous territory. In this new set-up, agents exist to serve and act on human intent aside from just emulating it. With this new ability to orchestrate and utilize functionality beyond its own textual skills, AI is able to produce a higher-value outcome for people and organizations through the medium of agents.

The opacity of chat-based AI naturally fails to lend itself to this new model. For agents to be truly autonomous, they need to be able to accurately interpret human intent and act on it—where there is consequence, therein lies the need for consent.

In agentic structures, that consent must be continuous and interruptible for it to hold its shape. This can't be granted without consistent oversight, without an ongoing dialogue between human and machine. Agentic systems must be transparent by design.

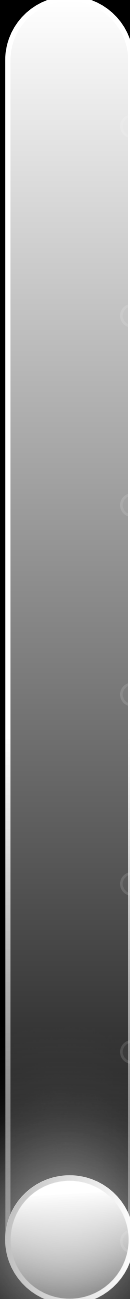
The business value of AI is something to be interrogated and not assumed. Organizations that are able to use discernment and recalibrate will see the most value from this next phase, while keeping their integrity intact.



If we want to truly integrate agentic AI into workflows, part of this is breaking away from the linear, black-box interface of chatbots, but it's also designing inherently transparent systems that we can predict and trust to work unanimously with humans.



Dominik Witzke
UX Design Lead, Star



Chapter one:

The next evolution in experience architecture

Experience design should be so seamless that it is practically invisible. Especially as agents become commonplace, building values into agentic systems becomes a prerequisite for trust, of which UX designers are the architects.

Empathic design is critical to building systems that are natively trustworthy, as well as protecting the identity of the user. While traditional UX principles still hold weight, they don't account for the nuances and challenges posed by AI agents. As agentic systems grow in capability, so will the expectation for organizations to define and divulge their own principles for agentic UX.

This is already being spearheaded by companies such as Microsoft, who have shared their new UX design principles for agents, and Anthropic who have published information about their upgraded agentic model capabilities, with Claude Opus 4.6 among the most recent examples at the time of writing.

These developments present a new onus for UX designers to shape the foundation for human-AI interactions.



From UX design to relationship design

For human intent to be shared accurately with AI agents, communication must be as close to faultless as possible. This type of interaction doesn't work with traditional chat-based interfaces; conversation history can get lost, which doesn't bode well for systems that rely on contextual company data.

In chatbot interactions, the experience is consistently user-driven. The AI receives instruction from the user and acts accordingly. But agentic capability goes beyond this model: agents can act on their own initiative and are able to make decisions and recommendations without a user present. So long as they have an objective, agents decide which steps to take to get there.

This fundamentally changes the role of the user, thus fundamentally changing the role of the designer. They control the environment in which human-AI interactions happen, and so they control how human intent is shared with the agent. Instead of simply designing interfaces, they start to design relationships.

Autonomy isn't binary, so consent shouldn't be treated as such either. To ensure agents interpret human intent accurately, exchanges between users and agents need to be inherently open. In this dynamic, consent is mutable and agent autonomy can be tailored to different levels of risk and preference.

Inclusive agentic UX design

When designing agentic systems, UX designers are responsible for ensuring the system and user journey are fair, which sets the tone for interactions between the AI and user while ensuring they do not unintentionally reinforce personal bias. For any user experience to be inclusive, it must be reflective of the foundation it's built on.

Ethical agentic UX design should incorporate fair user flows and default settings, inclusive data representation, transparency around how decisions or recommendations are made, mechanisms for human oversight, and enabled user feedback and reporting. Systems should also be tested among diverse groups to identify and prevent biased results.

Inclusive design can't wait until after the fact—if these aspects are not embedded at the start, agents can produce discriminatory or unfair results, damage user trust, limit access to services for certain populations, and expose companies to regulatory risks under frameworks such as the [EU AI Act](#), GDPR or consumer protection laws.





**Agents don't just respond; they act.
They take initiative and make decisions.
The challenge then becomes keeping humans
responsible even though the AI is acting.**

Dominik Witzke
UX Design Lead, Star

See-through strategy

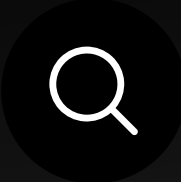
In chat-based interaction models, opacity is a given. The user types a prompt and receives a response, with little knowledge about what happened in between. For agents to be trusted to only do what they’re supposed to do, this formula has to change, and it has to be inspectable.

However, the agent-to-user dynamic also implies reduced image control; in a brand context, for example, when an agent is asked a query about a product, it responds with an interpretation, not direct communication from the brand. This leaves room for breakdown in brand messaging and subsequent risk to reputation.

It has implications on the agent side, too: in the event that an agent-driven product doesn’t fulfill expectations, and especially if it violates ethical or moral principles, trust can collapse quickly, and smaller firms may not survive the impact.

Cultivating trust through agentic systems is crucial to ensuring humans remain active orchestrators of agentic AI rather than passive participants. This demands an interaction strategy, one underpinned by transparency, accountability, bias awareness and clarity of ownership.

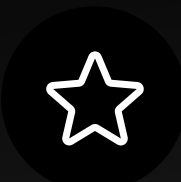
Steps for cultivating transparency through agentic UX



Contextual overlays:
The ability to inspect why a recommendation was made or action was taken (e.g. hovering over a system suggestion reveals which inputs/variables/historical signals influenced it).



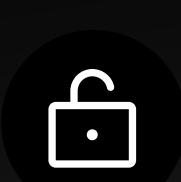
Reasoning summaries:
Accessible system explanations for autonomous decisions (e.g. “this action was triggered by...”)



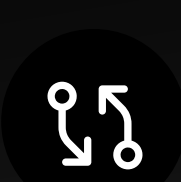
Confidence indicators:
Visual signals that indicate levels of certainty, probability and confidence.



In-line decision trails:
A displayed chain of actions showing how an objective was met and the steps taken to arrive there.

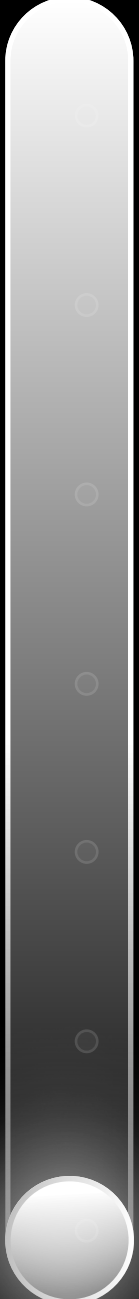


Direct information access without agentic interpretation:
The user has access to underlying raw data and information without any interpretation of AI in between to be able to check what the decision-making base was.



Canvas-based orchestration:
Open, node-based communication between human and machine where intent, constraints and output are visible.





Chapter two:

Modelling human- agentic interaction

Worth flagging

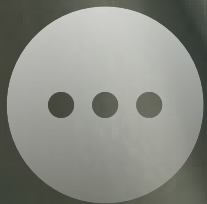
A 12% margin drop in this segment. Want to look closer?

Dismiss

Action Required

You've drafted a response to the project delay. How would you like to proceed?

Cross-referencing with Q4 benchmarks...



While AI-forward companies focus heavily on what the technology is capable of—its speed, its automation potential, very few are asking questions about how it and humans should interact, or how to safeguard human agency as these systems become increasingly powerful and, more recently, as they become autonomous. Star’s Human-Agentic Interaction (HAI) Model addresses this critical gap.

The HAI Model is based on the belief that as AI systems become more advanced, a shared set of principles are needed for organizations that create and operate them: designers, engineers, product managers, data scientists and business leaders. It provides guidance as to ensuring every AI-driven experience delivers real value to real people, respects the ethical and regulatory context it exists within, and keeps humans in control.



Core principles for human-agentic interaction

Human values come first: The agentic system's core purpose is to serve the user and create a clear, felt value-add (i.e. time saved, reduced labor, improved outcomes, reduced risk, entertainment, education); it should default to assistive guidance only, if it cannot demonstrate said value.

Safety over capability: Safety must be proactive and actions bounded by explicit permissions, risk level and context. The agentic system gracefully resorts to recommendations when confidence, context or authorization is insufficient.

Delegation must be revocable: Delegation to the agentic system must be explicit, granular and instantly revocable. Agents should never adjust or upgrade intent without permission: if asked to plan, it does not purchase; if asked to recommend, it does not execute.

Adaptive collaboration: Agents should be constantly learning and adjusting based on interactions with the user or other systems for optimized collaboration. Adaptation should only ever improve usefulness, and not manipulate or erode control.

Transparency through traceability: At any time, users must be able to understand and trace: what the system did or plans to do, the data and reasoning it used, the systems and tools it used, the permissions and credentials exercised, and what was delegated vs. user-approved.

Future-proof data and output: All outputs by the system must be stored somewhere accessible to the user and, if the agent is replaced or discontinued, the user retains full access to everything it created.

The system should actively support the possibility that the user may need such data independently in the future.

Architecturally embedded, not retrofit: Transparency, control, traceability and data sovereignty must be built in from inception, not bolted on. A system not designed with these principles from the start will face significant architectural constraints in fully honoring them.

Contextual compliance: Agentic systems must be ethical and compliant with the applicable standards and regulations in that context—jurisdiction, culture, domain. In the event of conflict, the system must clearly flag and explain the constraint, while offering safe alternatives.

These principles do not describe a single product or platform. They describe the conditions under which any agentic system, regardless of provider or architecture, can be trusted to act on behalf of a human. Organizations that embed them from the start will be better positioned to scale responsibly. Those that treat them as optional will find them impossible to retrofit when scale demands it.



Agent-to-Agent (A2A) Commerce

Such principles should apply across all sectors and disciplines in which agentic systems operate. The commerce sector is one example where agentic AI is already upending traditional infrastructure and the need for stronger governance and monitoring mechanisms is emerging.

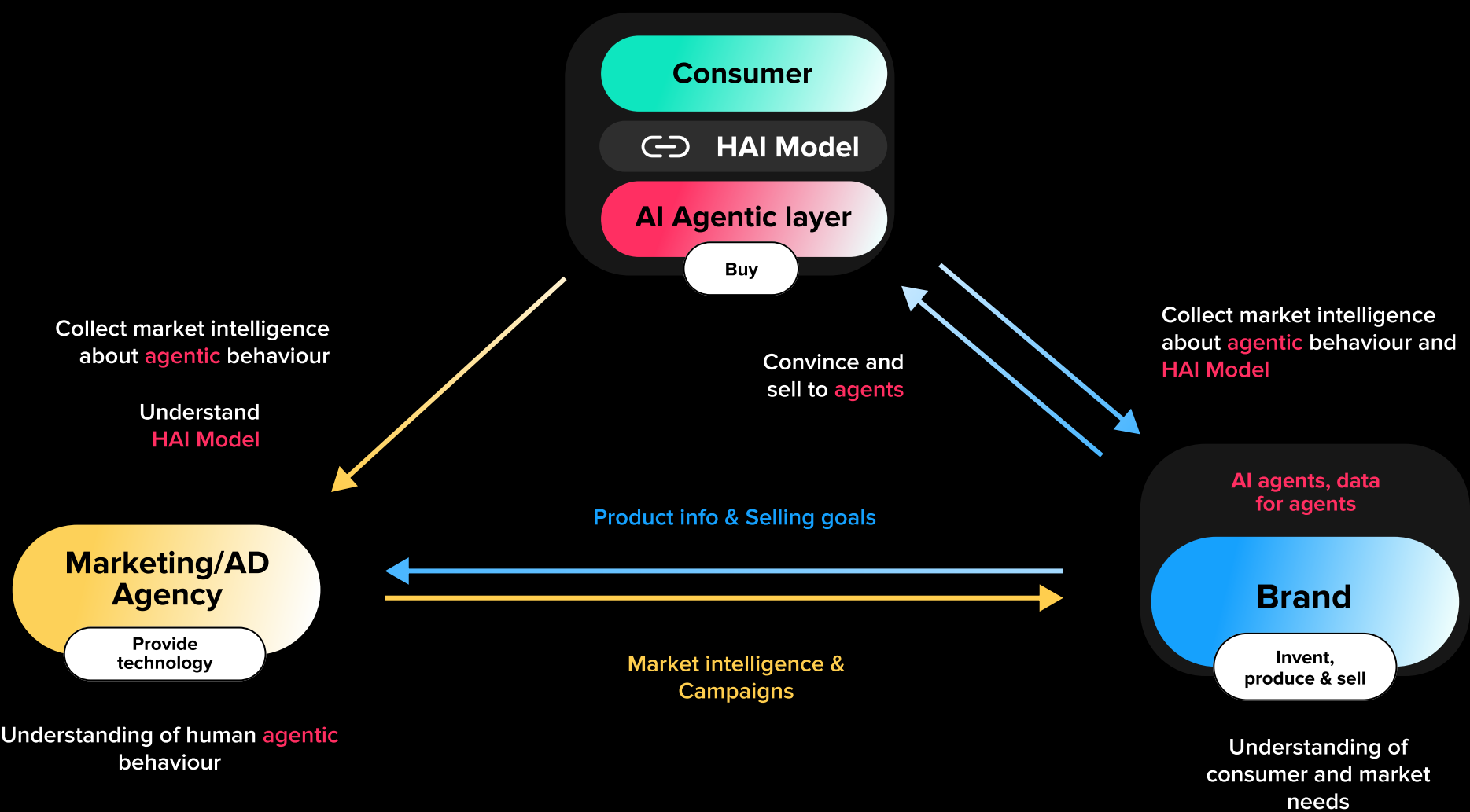
As virtual personal shoppers, AI agents have become the new mediators of brand perception, which means brands are no longer speaking directly to their customers. In this new world of commerce, where customers' digital proxies navigate the ecosystem, merchants operating online now have a third party to appeal to unless they start to provide agents of their own which can act as conduits for brand messaging and appropriately converse with user-proxy agents.



If you ask an LLM about a product, the response won't be messaging from the brand—it will be an interpretation of it. Brands have to figure out how to get beyond this boundary.



Martin Fix
Technology Director, Star



Since agents interpret, summarize and prioritize product data independently, they may not always synthesize information in a way that is correct or truly reflective of the source it came from. This puts advertising, representation, and intellectual property use all at risk if AI-generated descriptions or recommendations do not accurately reflect a brand's claims, and raises questions around liability when it comes to AI-driven recommendations and compliance with consumer and data privacy laws.

A2A commerce changes the traditional shopping flow altogether, reshuffling stages of the customer journey and prompting brands to deploy their own agents to preserve brand sovereignty. Without a brain, agents can't make decisions based on emotionally resonant strategies like brand storytelling, leaving consumer activity to be shaped by preferences, contextual memory and structured product data.

If brands want to win, they'll have to rely on substantial metrics like product data, trust signals, and alignment with user and brand intent in a way that ensures accurate interpretation of the intended message.



The three interaction models of agentic commerce

AI agents are coming to commerce in three ways:



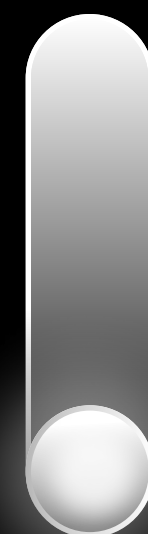
Agent to site

Agents interact directly with the brand or merchant to fulfill bookings and purchases.



Agent to agent

Agents autonomously converse and transact with other agents, e.g. a personal shopper agent communicates with an in-house commerce agent to negotiate a discount.



Brokered agent to site:

Intermediary AI systems allow for multiple agents to interact across multiple platforms, e.g. a vendor agent communicates with the broker agent of a restaurant-booking platform to find you a table and apply any relevant deals.



Toy vs. tool: evaluating AI readiness

Universal access does not equal universal readiness. Even as AI becomes more advanced, organizations should be careful to not overestimate their digital maturity and internal culture fit. When AI lacks cohesive governance and structure, that’s when adoption can go south: Gartner projects that by 2027, 60% of organizations will fail to realize the value they anticipated from AI because of this.

If an organization has established readiness, they should start small: practical experience in the form of pilot projects and small integrations can help to familiarize teams with AI, while realistically gauging the business impact. Experimentation should be about preparing to fail and learn, rather than trying to over engineer or optimize too early.

In the warm glow of mass adoption and productivity-filled promises, companies must stay grounded and use successful pilots to inform systemic structure and standardize processes to avoid fragmented adoption.

This approach also supports adoption from a cultural perspective: when pilots are successful, they breed confidence in the technology itself as well as the potential for it to positively impact existing processes and the organization as a whole.

Considerations for ethical AI adoption



Digital enablement:

In sectors where everything is done manually or by humans, AI has no leverage. Consider: Are your services already digitally enabled to some extent?



Availability of data:

For AI to be useful, it needs usable and accessible company-specific data. Without this, AI is superficial and unable to add real value.



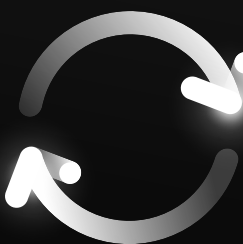
AI literacy:

It is not enough for companies to go in blind. When there is no internal understanding of what AI can do for your business, companies risk overpromising or misapplying. Outside expertise is invaluable.



Technical infrastructure:

Even if you understand AI and have existing digital processes, can you handle AI technically? Is it scalable? Can you measure it? If not, pilots could fail.



Change capacity:

AI transformation is largely cultural. If you cannot transform your company, if employees are resistant, it doesn’t make sense to start the journey—cultural readiness must outweigh technical readiness.



Progress, not perfection

When it comes to measuring innovative success, we rely on being able to prove that a piece of technology will:
a) only do what it's supposed to do and b) won't do anything wrong. With AI, such technical proof is extremely challenging. Sometimes even impossible.

In the early phases of AI adoption, these markers of meaningful AI use are more valuable than precise metrics.



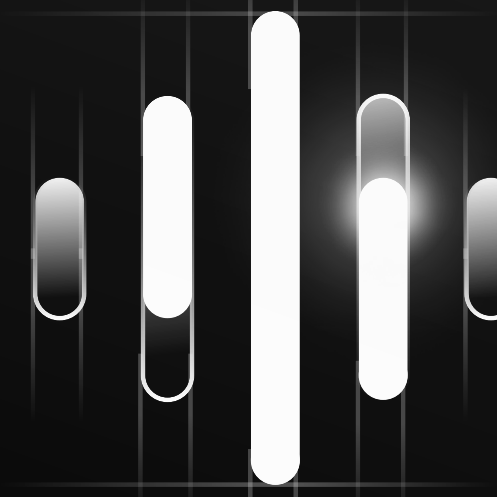
Human perception:

Ask employees about the perception of value and benefits gleaned from AI-driven processes.



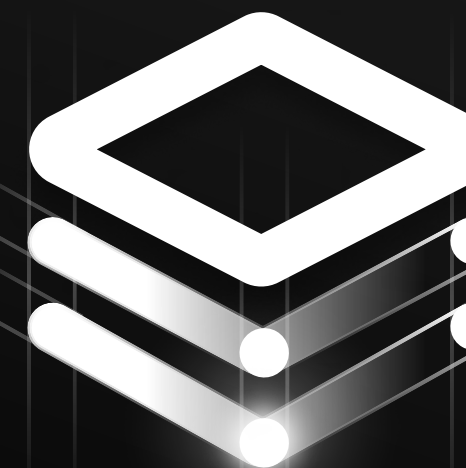
Directional trends:

Consider whether employees are more confident, more productive, more supported. Focus on overall sentiment rather than definitive numbers.



Count of AI-enabled processes:

One of the most practical ways to track progress is by the number of processes that are supported or fully run by AI. This alone can be an indication of meaningful integration and growing usage.



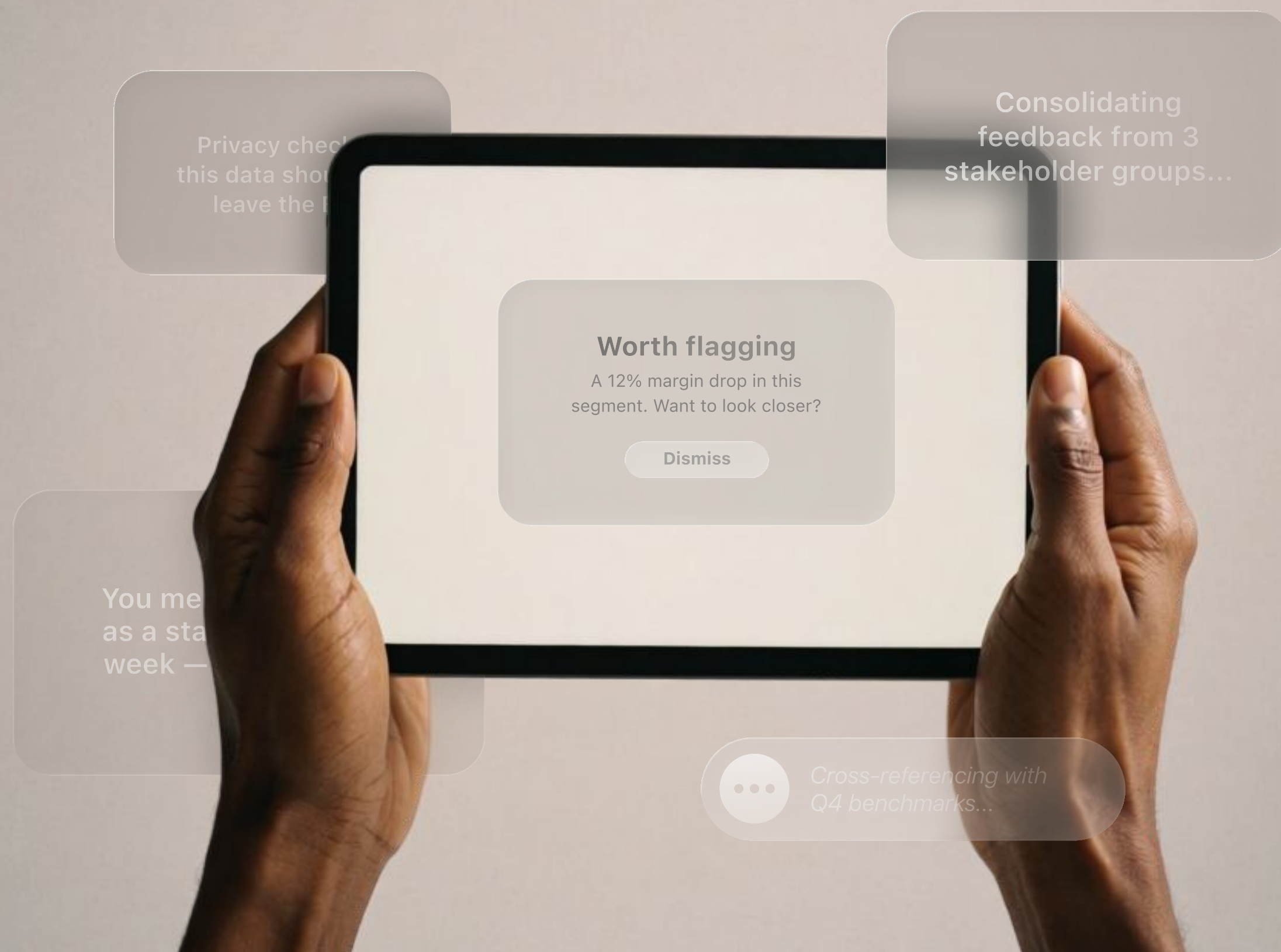
Financial indicators:

Where plausible, some correlations may point to general improvements, but don't rely on hard before-and-after baselines. Instead ask: Are costs gradually improving? Are margins responding over time?



The organizations that will get the most from AI are the ones that treat readiness as a serious precondition. They'll invest in data infrastructure before deploying models, develop internal literacy before scaling externally, and let pilots generate learning, before using that to inform structure. Speed of adoption shouldn't outweigh depth of preparation. A balance of both is required to ensure success.

Readiness, in this sense, is a continuous investment, in people, governance and the cultural conditions that allow AI to operate as infrastructure rather than as an ongoing experiment. The organizations building that foundation now are the ones that will be hardest to compete with later.



“

We understand how modern AI systems are built and trained, but specific outputs can still be difficult to explain or predict in specific situations and novel contexts. That is why measurement needs to combine objective task performance with human-centered outcome assessment on trust, usefulness, comprehension and learning.



Martin Fix
Technology Director, Star



Conclusion:

**The future of
agentic technology
is experience-led**

Agentic AI poses a new boundary around users which drastically alters the dynamic between user and machine. In the most radical version of this new interaction model, agents communicate exclusively with other agents. Human decision making may be even partially delegated, and the concept of brand reduced to a set of data points.

What we know now for certain, is that whether agent to employee or agent to customer, quality of data, transparency and continuous consent will play a crucial part in how this next chapter evolves—and this will all be dictated by the user experience. To see this through, organizations must be ready to change culturally and willing to make decisions based on use-case integrity; the case for AI must be justified and embraced internally for transformation to materialize.

AI viability is not a given. Deployment should always be objective-driven and positioned around the potential to genuinely improve business outcomes, not hollowed-out metrics. When AI is deployed without purpose, there is a human cost; when deployed with integrity, it compounds value.

Star supports organizations on that journey, bridging business strategy with experience-led adoption to create real value by:

-  Uncovering unmet user needs through behavioral and contextual insights.
-  Ensuring emotional consistency across all touchpoints by integrating brands from inception.
-  Developing frictionless user journeys that unlock revenue potential through service design.
-  Delivering scalable solutions validated by the market, through rapid prototyping, testing, and progression into product-level solutions.

Looking for a partner to help you start or scale your AI journey?

STAR.GLOBAL.COM

LET’S TALK LASTING IMPACT 

Contributors



Martin Fix
Technology Director

Martin is a senior technology leader with 25 years' experience in software development, machine learning/AI, technology management and research. He brings 15 years' experience in change management and organization building, acting as a mentor and coach for growing and experienced business leaders.

mfix@star.global



Dominik Witzke
UX Design Lead

With more than a decade across innovation and consulting agencies, Dominik specializes in human-centered user experience, with a focus on interaction and service design. He leverages design thinking, user research, insights and strategy to accelerate product development for Star's clients.

dwitzke@star.global



Antonina Burlachenko
Head of Quality and
Regulatory Consulting

Antonina brings a wealth of expertise spanning medical device regulation, software development, quality assurance and product management. As a certified lead auditor for ISO 27001 and ISO 13485, she oversees quality management, information security management, risk management (ISO 14971) and compliance at Star.

aburlachenko@star.global

About Star

Star is a global technology consultancy that helps enterprises bridge strategy and execution through deep industry expertise and an AI-native delivery engine to design, build and scale digital solutions with lasting business value and end-user impact.

1200+ projects

16+ years on market

450+ clients

Sunnyvale

San Francisco

Medellín

London

Copenhagen

Frankfurt

Munich

Wroclaw

Kyiv

Shanghai

Tokyo

Ho Chi Minh City