

**SuccessLab**

Three Decisions That Will Define the AI-Enabled Service Organization

Field notes from the 7th SuccessLab Executive Forum — London

By Omid Razavi, Ph.D. · Founder, SuccessLab

Customer service has proven AI can summarize, draft, route, and resolve. The harder work is making it safe, economical, and effective at scale. Drawing on the 7th SuccessLab Executive Forum in London, this piece lays out the three decisions that will define the AI-enabled service organization: how much freedom to give a customer-facing agent, how to prove value beyond productivity, and what operating model lets people supervise a growing AI workforce.

Customer service has cleared the first AI hurdle.

Organizations have shown that AI can summarize calls, draft replies, search knowledge, route work, and resolve selected inquiries. The harder challenge is turning those capabilities into a service model that customers trust, that delivers sustainable economics, and that operates at scale.

Three decisions will shape that model:

- 1. How much freedom should a customer-facing AI agent have?** Too little creates rigidity. Too much creates risk. The goal is useful freedom within clear boundaries.
- 2. How should organizations demonstrate value beyond productivity?** Containment alone does not prove value. The stronger case connects AI to journey economics, retention, customer lifetime value, and demand removed at the source.

- 3. What operating model lets people supervise a growing AI workforce?** Handoffs are becoming oversight. That requires clear ownership, grounded knowledge, governance, and an operating model built for life after launch.

These decisions are interconnected. Together, they determine whether AI improves customer outcomes and business performance or merely accelerates the operation already in place.

We built Customer Service Unbound: From Automation to Reinvention around these questions. Hosted by BCG, the 7th SuccessLab Executive Forum — London brought together senior leaders from banking, insurance, retail, energy, media, telecommunications, utilities, technology, and professional services to test these decisions against current practice.

The session was held under the Chatham House Rule. The observations and quotations that follow are intentionally unattributed.

AI is redesigning the service system

AI adoption is often described as a single transformation. The reality is broader. *“AI in customer service is hundreds of capabilities, decisions, and operating changes.”* Within the same organization, AI may draft emails, summarize calls, search knowledge, detect intent, complete workflows, and communicate directly with customers. Each creates value differently and carries distinct risks.

Across that variety, a consistent direction is emerging. Unstructured interactions are becoming business intelligence. Reactive support is moving toward proactive engagement. Channel deflection is giving way to resolution where the customer already is. Human roles are shifting from executing every interaction toward supervising more AI-led work. The scorecard is expanding beyond efficiency to include trust, customer outcomes, and business value.

That shift elevates the role of customer service. Customer interactions contain signals about product failures, operational gaps, unmet needs, retention risk, and commercial opportunity. AI can interpret those signals at scale and direct them to the teams positioned to act.

“AI’s real return extends far beyond productivity.”

The greater opportunity is to remove avoidable causes of customer contact and friction, and use service intelligence to improve the customer and business systems as a whole.

DECISION ONE

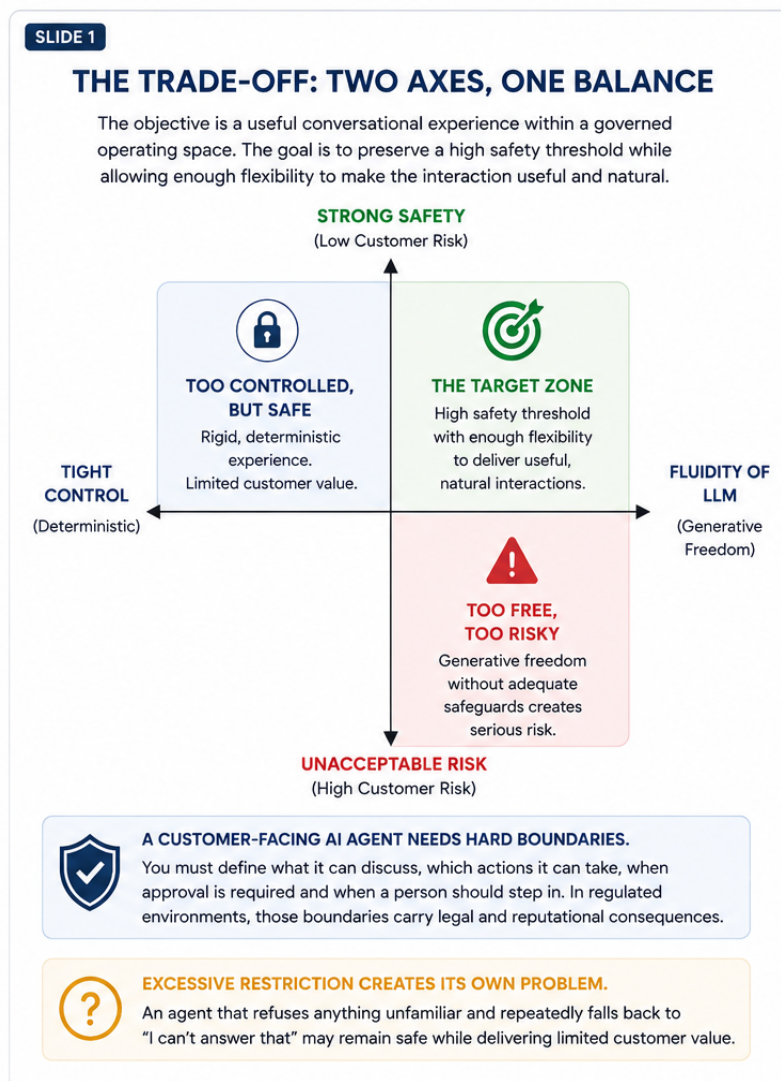
Balancing agent safety and customer value

The first decision concerns freedom.

A customer-facing AI agent needs hard boundaries. You must define what it can discuss, which actions it can take, when approval is required, and when a person should step in. In regulated environments, those boundaries carry legal and reputational consequences.

Excessive restriction creates its own problem. An agent that rejects anything unfamiliar and repeatedly resorts to “I can't answer that” may remain safe while delivering limited customer value.

The objective is a useful conversational experience within a governed operating space.



Two axes made the trade-off visible. One ran from tight control to the fluidity of a large language model. The other ran from strong safety to unacceptable customer risk.

Tight control and strong safety can produce a rigid, deterministic experience. Generative freedom without adequate safeguards poses a greater risk. The goal is to preserve a high safety threshold while allowing enough flexibility to make the interaction useful and natural.

Scope is where this becomes concrete. The instinct is to treat scope as binary: in or out. The reality has tiers.

Some topics sit at the center of the agent's mandate. Others sit plainly outside it, and a clear decline is appropriate. The harder case is the adjacent question, the one the agent was not designed to answer, but that a customer may reasonably ask.

A flat refusal can make the agent appear broken. A well-designed response can preserve trust by acknowledging the question, explaining the limitation, and offering a useful next step. Designing this middle ground requires considerable judgment.

SLIDE 2

SCOPE: NOT BINARY, BUT TIERS

The instinct is to think in binary terms: in scope or out.
The reality has tiers.

	1. CORE IN SCOPE At the center of the agent's mandate. It should handle confidently and well.	Example: Checking an account balance, payment due date, order status.
	2. ADJACENT / GRAY ZONE Related but not core. Customers may reasonably ask. Requires careful handling.	Example: Explaining options outside the existing product, or providing general guidance adjacent to the core mandate.
	3. OUT OF SCOPE Outside the agent's knowledge, permissions or purpose. The agent should decline and redirect.	Example: Tax advice, legal counsel, or actions the agent is not authorized to take.

**HOW THE AGENT HANDLES THE GRAY ZONE IS CRITICAL.**
A flat refusal makes it look broken. A graceful response preserves trust by acknowledging the question, explaining the limitation and offering a useful next step.

Safety must be designed in

One principle ran through the discussion: safety must be designed in. *“Safety is the foundation, not the final checkpoint.”* Safety cannot be applied only at the edges of an AI system. It must run through use-case selection, data permissions, knowledge, prompt and orchestration design, evaluation, monitoring, human review, and customer remediation.

The area many teams underdevelop is what happens after an incorrect answer reaches a customer. Who detects it? Who assesses the potential harm? Who contacts the customer and puts the situation right? Who corrects the system? Who remains accountable? Those responsibilities must be clear before the AI agent operates at scale.

Knowledge is decision infrastructure

A foundation model can produce a confident answer based on its training. Confidence, however, is not evidence.

The enterprise needs an authoritative source against which the answer can be evaluated. When the response and its validation both rely on the same ungrounded model, little has been independently verified. In one implementation discussed during the session, the immediate solution was to ground the agent in curated, human-reviewed knowledge.

“An AI agent is only as reliable as the knowledge grounding its decisions.”

That requirement exposes familiar enterprise problems: fragmented sources, conflicting versions, unclear ownership, stale content, and documentation written for employees rather than customers. AI is elevating knowledge management from a support discipline to a core part of enterprise decision-making. It determines what the agent knows, which actions it can justify, and how confidently the organization can defend the outcome.

Benchmark AI against reality

AI will make mistakes. Human service operations already contain errors, inconsistencies, and variations. *“AI should be measured against the operation it is expected to improve.”*

The responsible approach is to establish a credible human baseline for accuracy, compliance, consistency, resolution quality, rework, and customer effort. Many organizations do not yet have one. Without that baseline, leaders cannot make informed decisions about where AI may act independently, where it should assist, and where judgment should remain with a person.

Build on enterprise advantage

Enterprise AI also requires a realistic view of differentiation. *“Do not compete with foundation models on intelligence. Win with the context, trust, and authority only your enterprise can provide.”*

A company's advantage lies in proprietary customer context, trusted relationships, access to operational systems, industry expertise, and the authority to act. Those assets allow the organization to move beyond generating an answer. They enable it to understand the customer's situation, apply the right policy, and complete the next step. That is the foundation of a differentiated enterprise AI experience.

DECISION TWO

Proving value across the whole journey

If the first decision earns the right to deploy, the second earns the right to scale.

Most AI business cases begin with productivity: shorter handling times, higher containment, and lower administrative effort. Those measures are useful, but incomplete. A credible value case must also account for the total cost of operating the capability, its effect on customer behavior across channels, and whether the improved service contributes to retention and customer lifetime value.

Start with a bounded problem and scale with evidence

The strongest early deployments are usually bounded. Start with a narrow, well-understood category of demand. Allow the AI agent to assess each request before deciding whether to respond. Keep complex, sensitive, and emotionally consequential interactions with people until the evidence supports greater autonomy.

Email can serve as a useful starting point, as the system can review the request before determining whether it falls within its approved scope. The governing test is simple. *“AI creates value only when it improves the customer outcome.”*

Measure the whole system

The economics extend well beyond model cost. The full operating cost may include integration, data preparation, knowledge development, evaluation, monitoring, governance, human oversight, remediation, and ongoing product management. For some interactions, a person may still be the better economic choice, particularly where volume is low, complexity is high, or dependable automation costs more than it returns.

Containment also requires careful interpretation. An agent may resolve a large share of the interactions it receives while reducing total human effort by considerably less. Customers switch channels, return with follow-up questions, and create downstream work. The business case should therefore measure total customer and employee effort across the journey rather than relying on the headline containment rate.

The harder question is how service improvement connects to customer lifetime value. Efficiency is easier to count than trust and loyalty. That should raise the quality of the measurement, not narrow the ambition. Service AI can influence customer value through faster resolution, lower effort, earlier intervention, and stronger recovery after failure. As routine interactions become more automated, the moments that require human reassurance, judgment, and relationship-building may become more valuable.

Decide how the gains will be used

Then comes the question many programs postpone: what should happen to the capacity AI creates? The capacity can be captured as savings, reinvested in higher-value customer work, or divided between the two. Reinvestment may support growth, address more complex problems, expand proactive engagement, or strengthen customer relationships.

That decision should be explicit and made jointly with Finance. Customer leaders seeking reinvestment need to show in advance how it will create more value than capturing the capacity entirely as savings.

DECISION THREE

Designing the operating model for human oversight

Value at scale requires an operating model capable of governing and supervising AI-led work.

From handoff to oversight

The established model assigns routine interactions to automation and transfers difficult cases to a person. The emerging model keeps the AI agent active across more of the journey while a person supervises, approves, and intervenes.

In a control-tower environment, one employee may oversee several AI-led conversations. Risk signals may prompt them to approve a sensitive action, redirect the agent, intervene directly, or flag an interaction for review. This reduces unnecessary handoffs and the friction they create. It also changes the role of the service professional, from handling every interaction to supervising outcomes, managing exceptions, and improving the system.

Define ownership after launch

An AI agent needs a clear owner once it enters production. Models change. Knowledge ages. Customer behavior shifts. Policies, costs, and vendors evolve. Someone must remain accountable for performance, evaluation, knowledge maintenance, anomaly investigation,

compliance, and the approval of new capabilities. A project team can launch an agent. A permanent operating model keeps it reliable, up to date, and accountable.

Build the organization around the work

Human oversight requires new capabilities across service design, agent evaluation, knowledge engineering, governance, and exception management. Some capabilities will need to be brought in. Others should be developed internally. Frontline expertise remains essential because these employees understand customer language, operational nuance, and the moments that require judgment.

The work must also remain cross-functional. Product, design, engineering, data science, operations, knowledge, risk, and compliance all shape the outcome. *“AI will scale only as far as the operating model can support it.”* Progress improves when these disciplines make decisions together rather than passing them between functions.

Redesign before you automate

One principle underpins the operating model. *“Automating a flawed process usually scales the flaw.”* The journey should be examined before it is automated. Why does each step exist? Where does it create customer effort? Which controls remain necessary? Where does information break down? Which decisions require judgment? Which problems can be removed at the source?

The same scrutiny is changing sourcing decisions. The traditional onshore-or-offshore choice is becoming more granular. The relevant questions are which interactions should sit where, which require specialist judgment, which can be delivered effectively with AI augmentation, and which capabilities should remain within the organization.

Prepare for agent-to-agent service

Customers already use general-purpose AI to research products, test advice, and prepare for conversations with companies. The next stage is agent-to-agent service, where a customer's AI interacts directly with an enterprise AI. *“Soon, the customer journey will be shaped by two agents: theirs and yours.”*

That raises a new set of questions. How will an external agent prove whom it represents? What authority has the customer delegated to it? What information can it access, and which actions can it take? How will decisions be recorded, challenged, and reversed? Who is accountable when an agent acts incorrectly or the two systems conflict?

Safe agent-to-agent service will require trusted identity, explicit permissions, verifiable authority, observable actions, and clear accountability. Building those capabilities now will

strengthen today's customer experience and prepare the organization for the next service model.

What customer leaders must decide

The practical advice was clear. *"Start with a bounded problem, build the evidence, and scale with intent."* Each deployment should produce evidence, reusable capabilities, and greater confidence about where to scale next.

Leaders need a clear direction and shorter execution cycles. A rigid twelve-month roadmap is unlikely to survive unchanged. The next phase turns on three decisions: balance agent safety with customer value, connect AI economics to customer and business outcomes, and design effective human oversight of customer-facing AI.

Customer service sits at the intersection of trust, execution, product performance, and growth. AI increases both the opportunity and the responsibility of that position.

The technology will continue to improve. The quality of these decisions will determine whether AI produces isolated efficiency gains or strengthens the service model and the wider business.



ABOUT THE AUTHOR

Omid Razavi, Ph.D., is the founder of SuccessLab and the author of CCO Perspectives. He has spent more than thirty years in enterprise software, in customer success, support, and service leadership roles at companies including ServiceNow, SAP, and SupportLogic. Through SuccessLab, he convenes senior customer, revenue, and technology leaders across the United States and Europe to work through the decisions that shape the next operating model for customer-facing organizations.

SuccessLab is a community and advisory brand at the intersection of AI transformation and enterprise customer leadership. It runs the SuccessLab Executive Forums and Roundtables, publishes practitioner-level field notes and analysis under CCO Perspectives, and brings revenue and customer leaders together to learn from one another as AI reshapes how service is designed, delivered, and valued.

Subscribe to [CCO Perspectives](#) and join the community at community.successlab.us.