

Deploying conversational AI as voice agents is transforming customer service—especially in complex sectors like retail and telecom. Yet one persistent and costly problem remains: voice agents sometimes "invent" discounts or goodwill credits that don't exist, leading to frustrated customers and expensive post-call scrambles.

Industry leaders like **Suprmind**, **Air Canada**, and **OpenAI** have tackled this challenge head-on. Their approaches leverage a mix of *retrieval-augmented generation* (RAG) technology, robust *speech-to-text* and *text-to-speech* pipelines, and critical policy controls.

In this article, we'll explore the seven main failure points where voice agents go off the rails, the limits of RAG and knowledge base hygiene, why live tools must be treated as the source of truth, and how high-precision entity confirmation plus readback practices can make "invented discounts" a thing of the past.

Seven Failure Points in Voice Agents That Lead to "Invented" Discounts

Before understanding solutions, it's important to identify where voice agents break down:

1. **Unstable Knowledge Bases:** Stale or incorrect discount policies cause incorrect responses.
2. **RAG Overconfidence:** The agent "hallucinates" when retrieving ambiguous or incomplete data.
3. **Lack of Real-Time Data Integration:** Discounts tied to live promotions or customer history are inaccurately reported.
4. **Ambiguous Customer Inputs:** Misrecognitions in speech-to-text pipelines lead to errors in applying goodwill.
5. **Loose Entity Extraction:** Poorly tuned NLP models extract wrong amounts, dates, or policy references.
6. **No Clear Authority Control:** The agent lacks a strict "never promise" rule that forbids unauthorized goodwill grants.
7. **Absent Confirmation & Readback:** Customers are not asked to confirm understanding before finalizing offers.

Anyone managing voice AI deployments should maintain a notebook of real call snippets illustrating these failure modes—examples like "B three one seven two"—help connect theory to actual customer interactions.

RAG Limits and Why Knowledge Base Hygiene Matters

You know what's funny? retrieval-augmented generation (rag) is a game-changer in conversational ai, allowing voice agents to combine generative language models with external knowledge repositories. However, even the best RAG systems have inherent limits:

RAG Limitation	Description	Impact on Discounts/Goodwill	Source
Ambiguity	Retrieval relies on keyword matching and similarity measures, sometimes pulling irrelevant or outdated documents.	Agent cites invalid or obsolete discount policies.	Context Loss
Loss	The generative layer can fill in data gaps, leading to "hallucinations" when source docs lack specific details.	Confidently invents goodwill amounts.	Update Latency
Knowledge bases	Often update on daily or longer cycles, missing last-minute promotions or policy shifts.	Fails to reflect current applicable discounts.	

Maintaining exceptional knowledge base hygiene—regular curation, timestamping, verification against policy documents—is essential. But great hygiene isn't enough on its own. Systems need a live data layer anchored in transactional databases and policy management tools.

Live Tools: The Ultimate Source of Truth for Customer-Specific Facts

From Suprmind's recent implementations to Air Canada's agent assistive tech powered by OpenAI, a common pattern emerges: voice agents must pull live data from authoritative tools making real-time decisions.

- **Customer Account Systems:** For up-to-the-minute eligibility and goodwill credit records.
- **Policy Rules Engines:** That programmatically enforce "never promise" lists and discount ceilings.
- **Promotions Databases:** Reflecting current valid discounts and expiry dates.

Failing to integrate these live systems is tantamount to relying on guesswork. This is why *authority controls* at the API and middleware layer are essential to prevent the AI from issuing unauthorized promises.

Authority Controls and the "Never Promise" List

Two related concepts are your strongest weapons against rogue goodwill promises:

1. **Authority Controls:** These ensure the voice agent's responses are governed by explicit permission layers. For instance, an agent cannot generate a goodwill credit exceeding a preset limit without supervisor approval.
2. **Never Promise List:** This is a dynamically maintained set of offers, discounts, or credits the AI must categorically never mention or authorize independently.

I remember a project where wished they had known this beforehand.. Implementing these involves rigorous policy rules engines—capable of evaluating context, customer history, and regulatory constraints in real-time. It's best practice to externalize such [Helpful resources](#) engines rather than bake rules into GPT-style prompt instructions, which degrade over time.

High-Precision Entity Confirmation and Readback

Even with perfect data, speech recognition and entity extraction issues abound. Misheard numbers, misinterpreted policy codes, or jumbled customer names can derail accuracy.

Voice agents must adopt a robust confirmation and readback protocol:

- **Entity Extraction Thresholds:** Only consider entities valid if confidence scores cross a strict threshold.
- **Explicit Confirmation Steps:** Ask customer to confirm quoted discount amounts, reference numbers, or goodwill credits.
- **Repeat Back Key Details:** The agent paraphrases offers, dates, and values clearly before closing the interaction.

This reduces variance introduced during speech-to-text and safeguards the agent from unverified self-generated content.

Putting It All Together: Best Practices Summary

Challenge	Recommended Solution	Example Tools
RAG hallucination in discount claims	Strict knowledge base hygiene + live policy rules engines	Air Canada live discount APIs, Suprmind policy rules
Outdated or ambiguous KB data	Frequent content curation + retrieval score thresholds	OpenAI-based retrieval pipelines with filtering
Speech recognition errors on monetary values	High-precision entity extraction + confirmation prompts	Custom-built STT pipelines with confidence metrics
Unauthorized goodwill promises	Authority controls enforcing never promise lists	Policy rules engines integrated at API level
Customer confusion from unclear readbacks	Explicit readback protocol with customer verification	Voice UI design best practices

Final Thoughts: Don't Blame "Hallucinations" — Enforce Guardrails

Whenever you hear someone blame an AI failure on "hallucinations," I always ask: *What is the source of truth for that sentence?* Blaming hallucinations alone lets teams off the hook for proper systems design and controls. Promises of discounts or goodwill credits must not live only inside a prompt or language model output—it must be backed by authoritative live data and programmatic guardrails.



Deploying best practices from pioneers like Suprmind, Air Canada, and OpenAI, and leveraging tools such as RAG, speech-to-text, and text-to-speech pipelines intelligently will keep your voice agents truthful, trustworthy, and financially sound.



Remember: the voice agent is only as good as the integrations and authority controls built around it. Invest there first, and the "invented discount" woes will soon be behind you.