

Dewly AI Voice Agent

INBOUND & OUTBOUND CALLS FOR SHORT-STAY OPERATORS

Extending Dewly's AI duty manager from chat to the phone line: a voice agent mapped from the existing LangGraph agentic system, so guest calls and staff dispatch run on one reasoning core across two channels.



Scan for the live demo rmitvn-showcase.com/voiceagent

P50 ≤ 1.5s

turn latency, 800ms first audio

4 Domains Agents

pricing/availability, cancel/modify, escalation, search

12 quality metrics

experience, diagnostic & latency axes

One reasoning core, two channels: chat and voice. A fix to the duty-manager logic ships to both at once.

DEWLY AGENT WORKER

Real caller
phone / SIP / WebRTC

EVA
kind=AGENT worker — plays caller persona

audio in (mic)
reply audio out
publishes caller audio
subscribes reply + filler

LiveKit Room

shared transport only
— no logic lives here

caller audio
2 audio tracks

entrypoint(): RoomIO starts

STT → llm_node() → graph → synthesizer

Filler speech — runs in parallel
EMA latency estimate per domain

2 audio tracks → Room
reply + filler

RecorderIO
saves the input audio it received

Background & Motivation

Short-stay operators run on the phone: guests call to extend a stay, report a broken air-conditioner or ask for a late check-out, while one duty manager covers dozens of rooms. Dewly already answers these requests in chat, but the phone line stays manual calls are missed at night, context is lost when staff are dispatched, and every conversation is re-typed into the log. Voice carries the most urgency and the least automation, and it is also the hardest to judge: a wrong tone or a two-second silence breaks a call in a way a chat message never would.

Objectives

- Map the chat agent onto a real-time voice pipeline without forking the reasoning core.
- Handle both directions: inbound guest calls and outbound dispatch to housekeeping and maintenance.
- Hold a P50 of 1.5 s per turn, with first audio inside 800 ms.
- Build a repeatable evaluation framework so voice quality is measured per turn, not judged by ear.
- Stop within 300ms on a real interruption, without cutting off natural backchannels

Tech Stack

LangGraph LiveKit Gemini Deepgram
Cartesia Twilio telephony EVA framework

Methodology

- Channel & requirement mapping**
Audited the chat agent's tools and intents against real call recordings to scope what voice must support.
- Porting the LangGraph core**
Wrapped the graph behind a transport-agnostic interface so chat and voice share one set of nodes, tools and prompts.
- Voice pipeline integration**
Connected Twilio telephony to a LiveKit agent worker with Deepgram STT and Cartesia TTS, tuning turn detection and barge-in.
- EVA evaluation loop**
Scored every prototype call on responsiveness, speakability and tone, then fed failing turns back into prompt and pipeline changes.
- Pilot & deployment**
Ran the agent against hotel test lines and human callers for 60 days, logging every room for review and regression.

Conclusion & Future Work

Routing voice through the existing LangGraph core proved the cheaper path: the duty manager gained a phone line without a second agent to maintain, and EVA turned "does it sound right" into a number the team could act on each iteration. Next: sub-second P50 with speculative TTS, full-call rubrics such as task completion and escalation accuracy, and a link into the property management system so a call updates the booking directly.

Evaluation Framework - EVA

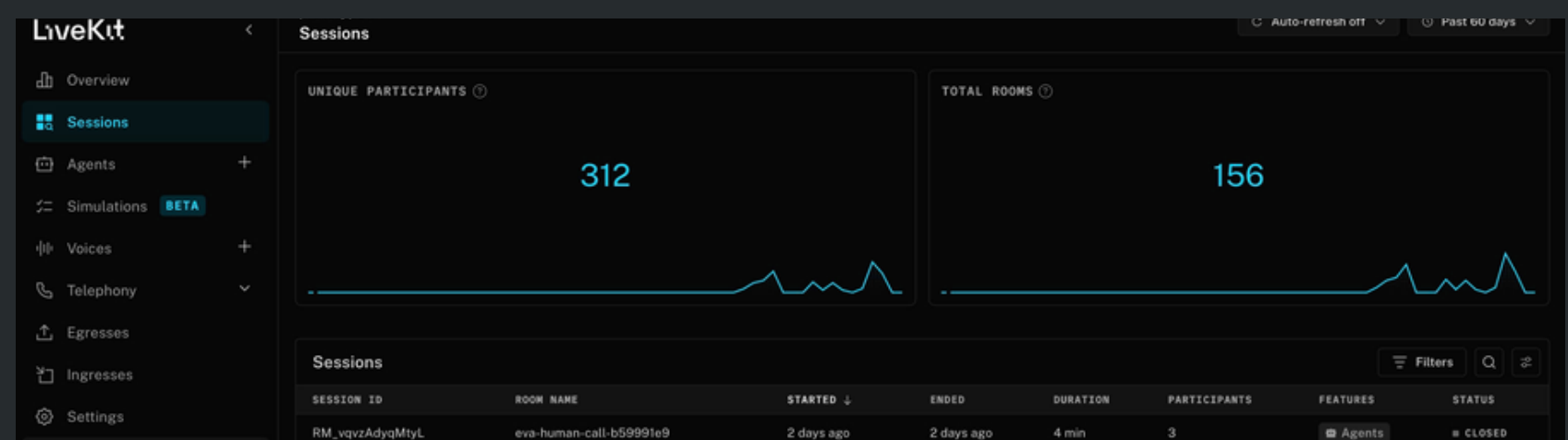
An LLM/audio/deterministic judge mix scores 12 metrics per turn - no composite score; each reports independently

```
1834 {
1835   "speakability": {
1836     "name": "speakability",
1837     "score": 0.857,
1838     "normalized_score": 0.857,
1839     "details": {
1840       "per_turn_ratings": {
1841         "1": 0,
1842         "2": 1,
1843         "3": 1,
1844         "4": 1,
1845         "5": 1,
1846         "6": 1,
1847         "7": 1,
1848         "8": 1,
1849         "9": 1,
1850         "10": 1,
1851         "11": 1,
1852         "12": 1
1853       },
1854       "per_turn_explanations": {
1855         "1": "The text contains markdown bullet points (*) and bold formatting (**863 Pierce Ave**), which violate the
1856         \"5\": \"The text is written in natural language without any markdown formatting, emojis, missing spaces, or unp
1857         \"6\": \"The text contains only standard punctuation and words. There are no visual formatting elements, emojis,
1858         \"7\": \"The response is plain text, easily spoken, and free of any markdown, emojis, or formatting violations.\"
1859         \"8\": \"The response consists of regular sentences with standard punctuation. It does not contain any formatin
1860         \"9\": \"The text includes standard punctuation and a number sign (#859), which is naturally spoken as 'number'.
1861         \"12\": \"The text is purely conversational and contains no visual formatting, markdown, emojis, or unpronounce
1862       },
1863       \"judge_prompt\": \"You are an expert evaluator analyzing whether text is voice-friendly and appropriate for text-
1864       \"judge_raw_response\": \"\"\"json\\n\\n {\\n  \\\"turn_id\\\": 1,\\n  \\\"explanation\\\": \\\"The text contains markdown
1865       \"aggregation\": \"mean\",
1866       \"num_turns\": 1,
1867       \"num_evaluated\": 7
1868     },
1869     \"error\": null,
1870     \"is_not_applicable\": false
1871   }
1872 }
```

EXPLANATION

Every call produces a structured, per-turn breakdown - a score and a written reason for each metric, not a single opaque number.

Results



LiveKit Sessions dashboard, showing an overview of recorded call sessions on the prototype environment.

Events	Participant	Event	Data	Timestamp
-	-	Room created	Event payload	20 Aug 2026 at 15:55:58
-	-	API	Event payload	20 Aug 2026 at 15:55:58
human-caller	-	Participant joining	Event payload	20 Aug 2026 at 15:56:10
human-caller	-	Participant active	Event payload	20 Aug 2026 at 15:56:10
eva-human-recorder	-	Participant joining	Event payload	20 Aug 2026 at 15:56:11
eva-human-recorder	-	Participant active	Event payload	20 Aug 2026 at 15:56:14
agent-AI_36c8gGao2hu	-	Participant joining	Event payload	20 Aug 2026 at 15:56:17
agent-AI_36c8gGao2hu	-	Participant active	Event payload	20 Aug 2026 at 15:56:18
human-caller	-	Participant left (CLIENT_INITIATED)	Event payload	20 Aug 2026 at 15:59:23
eva-human-recorder	-	Participant left (CLIENT_INITIATED)	Event payload	20 Aug 2026 at 15:59:25
agent-AI_36c8gGao2hu	-	Participant left (ROOM_DELETED)	Event payload	20 Aug 2026 at 15:59:27
-	-	Room ended	Event payload	20 Aug 2026 at 15:59:29

LiveKit session event log, showing the participant join/active/leave timeline for a single voice call.