

## WHY FINN

# Not a wrapper. The whole voice layer.

Most stacks chain providers in series and pay the latency tax at every hop. Finn runs one fused pipeline — voice intelligence, orchestration, observability, integration, and security — built for teams that ship at production scale.

&lt;400ms

End-to-end latency

46

Languages, live

99.99%

Uptime

50M+

Minutes shipped

## Four reasons teams pick Finn

- Fused pipeline. Streaming speech-to-text, inference, and text-to-speech in one hot pipeline — sub-400ms. Customers hear a real cadence, not stop-and-wait.
- No-code orchestration. Drag-and-drop call flows, conditional branching, human handoff, multi-agent coordination — orchestration that survives the messy real world.
- Built for regulated industries. Healthcare, finance, insurance, legal — from day one. SOC 2 Type II, HIPAA, GDPR, ISO 27001.
- Outcome observability. Recordings, transcripts, sentiment, and cost-per-outcome captured automatically. Spot drift, prove ROI, ship improvements in days not quarters.

## One agent vs. a room of dialers

| Human agent        | Finn                  |
|--------------------|-----------------------|
| 35–60 dials/hour   | 10,000+ concurrent    |
| 8–15% connect rate | 30–40% connect rate   |
| 15–20% talk time   | 90%+ talk time        |
| Labour + overhead  | Usage-based talk time |

Book a demo — [www.hirefinn.ai/bookademo](https://www.hirefinn.ai/bookademo)

Or start free at [hirefinn.ai](https://hirefinn.ai) · [hello@hirefinn.ai](mailto:hello@hirefinn.ai)