

1 - Как AI-агент входит во встречу (это дополнительный участник)

Приложение -> Бэкенд (комната + токен, диспатч) -> Сервер LiveKit (Cloud / self-host)
-> AgentServer запускает изолированное задание -> агент входит в комнату как участник.
Цикл агента: подписка на аудио -> STT -> LLM -> (опц.) TTS
-> публикация транскрипта / записи / речи обратно в комнату.
Диспатч < 150 мс. Каждое задание - свой процесс; живёт, пока все не выйдут.

2 - Три типа речевого конвейера (выбор по нужде функции)

STT -> LLM -> TTS : живые транскрипты, контроль. Дефолт для ноутейкера.
Realtime-модель : естественно, но БЕЗ промежуточного транскрипта - добавьте STT.
Half-cascade : realtime-понимание + отдельный голос TTS. Смесь обоих.

3 - Уровни LiveKit Cloud (2026 - перепроверьте на livekit.io/pricing)

Build \$0/мес	5 000 WebRTC, 1 000 агент, 50 ГБ, 5 одновр.
Ship \$50/мес	150k WebRTC (\$0.0005), 5k агент (\$0.01), 250 ГБ (\$0.12)
Scale \$500/мес	1.5M WebRTC (\$0.0004), 50k агент (\$0.01), 3 ТБ (\$0.10)
Enterprise	Индивидуальные цена и одновременность.
Self-host	Сервер + Agents = бесплатно (Apache 2.0); платите за инфру.
Важно	Модели (STT / LLM / TTS) тарифицируются поверх любого уровня.

4 - Сколько стоит минута (типичные тарифы моделей 2026)

Ноутейкер «слушать» = агент \$0.01 + STT \$0.006 + наблюдаемость \$0.01 ~ \$0.026 / мин
Говорящий агент = + LLM \$0.0077 + TTS \$0.030 ~ \$0.0635 / мин
Text-to-speech - самый тяжёлый слой; отказ от него примерно вдвое снижает стоимость.

5 - Чек-лист build-против-buy

- ☐ Сначала решите: ассистент - ваш продукт или рядовая функция транскрипта?
- ☐ Нужны живые транскрипты? Используйте STT-LLM-TTS или добавьте STT к realtime-модели.
- ☐ Ноутейкер «только слушать»? Уберите text-to-speech - стоимость минуты упадёт вдвое.
- ☐ Строгая приватность или место данных? Self-host open-source-сервера (Apache 2.0).
- ☐ Нужны транскрипты быстро? Купите сервис бота встреч (напр. Recall.ai ~\$0.50/час).