

Главный сдвиг: один фиксир. микс -> микс под слушателя

40 лет звук выходил как один фиксированный микс, одинаковый для всех. Следующее десятилетие хранит ингредиенты - диалог, музыку, эффекты - отдельно плюс инструкции, чтобы язык, акцент, уровень диалога и 3D-расстановка выбирались под слушателя при воспроизведении. Движок - обучаемые (нейросетевые) системы, что захватывают, чистят, сжимают и даже генерируют звук за один проход.

Что реально сейчас, а что обещание

Возможность	Что делает	Статус 2026	Подвох
Объектный контроль диалога	поднять речь, не музыку	в проде (ATSC 3.0/MPEG-H)	нужны объекты + ресивер
Обучаемое усиление речи	достать речь из микса	в проде (Dialog+, Amazon)	хуже на плотных миксах
Нейро-разделение источников	голос/музыка/эффекты	в проде; на устройстве	артефакты при наложении
Экспрессивный перевод речи	тот же голос, новый язык	ранний прод (Seamless)	задержка, точность, согласие
Конверсия акцента в реал. вр.	чётче, с эмоцией	запуск (Krisp 2026)	сарказм/срочность не решены
Визуальный дубляж (звук->видео)	губы под дубляж	в кино (Flawless)	цена, права, дипфейки
Сквозной обучаемый фронтэнд	одна сеть на всё	наука / частично	вычисления; Orus в транспорте
Полностью генеративный микс	модель собирает микс	лаборатория / концепт	вычисления, контроль, источник

Цифры, что делают это реальным

- Чёткий диалог - это мейнстрим: в полевых тестах Fraunhofer Dialog+ (WDR/BR) 83% зрителей оценили опцию, а 90% людей 60+ часто не разбирают речь в ТВ.
- Amazon Dialogue Boost: >86% тестеров предпочли усиленный диалог, 100% среди людей с потерей слуха; модель сжали до <1% размера для работы на устройстве.
- Математика громкости: музыка -24 LUFS против диалога -27 LUFS = голос на 3 LU ниже. Нормализация всего микса хранит разрыв; подъём только речи +6 LU - голос на 3 LU выше.
- Персонализация в масштабе: Netflix добавил 13 000+ часов аудиоописания на 34 языках в 2025 (+30% за год).
- Сжатие и генерация - один навык: кодек, что восстанавливает голос из ~1 кбит/с, в шаге от модели, что его выдумывает - поэтому происхождение крепится на уровне кодека.

Прежде чем закладывать дорожную карту - барьеры

- ☐ Читайте таблицу сверху вниз: верх = брать/строить сейчас; середина = пилот; низ = наблюдать, не закладывать.
- ☐ Каждую возможность из нижней половины гоните на ХУДШЕМ входе: шумный звонок, плотный микс, сложный акцент, слабое устройство.
- ☐ «Демо работает» - это не «продукт готов»: разрыв в длинном хвосте, который демо обошло.
- ☐ Для обычного живого транспорта Orus + настроенный jitter buffer всё ещё выигрывают по цене и надёжности.
- ☐ Любая функция с узнаваемым голосом:стройте согласие + раскрытие. EU AI Act ст. 50 - с 2 авг. 2026.
- ☐ Ставьте водяной знак на генерируемый/изменённый звук с первого дня - дооснащать происхождение дороже.