

Call for Participation

Overview

New in LRAC 2.0

Training data expands in size and language diversity; evaluation extends beyond **English (ENG)** to **Spanish (SPA)** and **Mandarin Chinese (CMN)**.

- **Hard cap of $\leq 1.5\text{M}$ parameters** — a firm eligibility constraint.
- **Theoretical memory-bandwidth penalty** applied to the ranking score.

- Quality scaling across the bitrate caps to transparency in clean speech
- Robustness to mild real-world distortions (background noise, reverberation)
- Behavior across ENG, SPA, and CMN

Bring new designs. We explicitly invite diversity beyond the Conv-Encoder + RVQ + Conv/GAN-Decoder pattern that dominated LRAC 1.0 — e.g., hybrid parametric + neural designs, classical-DSP front ends with neural quantizers/decoders, diffusion/flow decoders, scalar/lattice/product quantization, multi-rate architectures, and neural-augmented classical codecs.

A mandated, fixed public dataset — expanded to **~1,200 h** after curation (~2× LRA1.0), spanning **12 languages** with a roughly balanced ENG / non-ENG split. Sampling rate **24 kHz**. Training, validation, and tuning strictly on the designated challenge data.

A newly collected, withheld test set. Evaluation matches LRAC 1.0 Track 1 across five bitrate caps **{1, 6, 12, 24, 32} kbps**. Multilingual: clean-speech quality for SPA & CMN at higher rates; intelligibility for ENG, SPA, CMN. All evaluation languages appear in training.

Participants process open and withheld test sets and submit the stimuli. Both objective metrics and crowdsourced subjective tests are used; objective metrics are informational only. **Final rankings rest on subjective scores** — accounting for speech quality scaling toward transparency and speech intelligibility. **In Track 2**, the score is additionally adjusted by a theoretical memory-bandwidth penalty favoring leaner solutions at comparable quality and bitrate. The same battery below applies to both tracks.

Subjective evaluation battery — applies to both tracks. Each listed bitrate is a cap: an upper limit a system's constant rate may meet but not exceed (for example, 16 kbps satisfies the 32 kbps cap; 48 kbps does not). Clean-speech MUSHRA-1S assesses transparency against the hidden clean reference, with results for Opus at 32 kbps reported as a strong conventional-codec baseline. The framework accommodates full-battery testing of at least 7 submissions per track. If more are received, an objective-metric-based qualification round, detailed in the Official Rules, determines which submissions advance to the full evaluation. Rankings are based on a weighted battery of the subjective evaluation results. Track 1 is ranked on these subjective scores alone; Track 2 adds a memory-bandwidth penalty, favoring solutions with lower memory-bandwidth requirements at comparable quality and bitrate. The highest-ranked submission in each track is declared the winner. No financial prizes are awarded.

Join the LRAC 2.0 Challenge

Building practical speech codecs, researching novel neural architectures, or pushing low-resource coding toward transparency across languages? Register and help shape the next generation of low-resource communication systems.

Questions & Slack invite
lrac-challenge@cisco.com

