



LREC 2026

**Speech Language Models in Low-Resource Settings:
Performance, Evaluation, and Bias Analysis
(SPEAKABLE) @ LREC 2026**

Workshop Proceedings

Editors

**Nina Hosseini-Kivanani, Alessio Brutti, Marco
Matassoni, Sandipana Dowerah, Davide Liga**

11 May 2026
Palma, Mallorca (Spain)

©ELRA Language Resources Association (ELRA), 2026

These proceedings are licensed under a Creative Commons Attribution-
NonCommercial 4.0 International License (CC BY-NC 4.0)

ISBN 978-2-493814-83-8
EAN 9782493814838

Preface

Welcome to the SPEAKABLE 2026 Workshop on *Speech Language Models in Low-Resource Settings: Performance, Evaluation, and Bias Analysis*.

SPEAKABLE 2026 is a full-day workshop co-located with LREC 2026. It brings together researchers and practitioners working on speech-native language models, with a particular focus on low-resource languages, dialects, and speaker communities. The workshop was created in response to a growing need in the speech and language technology community: while speech foundation models and Speech LLMs have advanced rapidly, their benefits remain unevenly distributed across languages, domains, devices, and user groups.

Low-resource speech settings continue to face persistent constraints related to limited training data, uneven annotation quality, scarce evaluation benchmarks, and restricted computational budgets. These difficulties are often intensified by real-world deployment conditions, including accent and dialectal variation, channel and microphone mismatch, spontaneous speech phenomena, code-switching, noisy environments, and limited availability of lexicons or grapheme-to-phoneme resources. As a result, systems that appear strong under standard benchmark conditions may still fail to provide reliable, fair, and useful performance for many underrepresented communities.

The goal of SPEAKABLE 2026 is to provide a focused forum for addressing these challenges through three closely connected strands. The first strand is **efficient adaptation**, including parameter-efficient fine-tuning, multilingual transfer, knowledge distillation, and edge- or streaming-constrained inference for low-resource speech tasks. The second strand is **meaningful evaluation**, with an emphasis on moving beyond aggregate scores such as WER toward task-appropriate metrics, calibration, robustness analysis, abstention, and slice-aware reporting by accent, dialect, channel, speaker group, and speaking style. The third strand is **responsible practice**, treating bias analysis, data documentation, synthetic-data disclosure, privacy, and safety considerations as routine parts of scientific reporting rather than optional additions.

The call for papers welcomed work on efficient adaptation of Speech LLMs for low-resource languages, evaluation methods for ASR, spoken language understanding, and speech generation, robustness under domain shift, cascaded versus end-to-end error propagation, low-resource corpus creation, lexicon and G2P development, and ethics-by-default reporting. We especially encouraged submissions that combine methodological innovation with strong empirical evidence, transparent documentation, calibrated uncertainty, and, where possible, openly released resources or code.

This first edition of SPEAKABLE reflects the increasing importance of speech technologies for the long tail of languages and communities. The accepted contributions address a broad range of topics, including multilingual and cross-lingual modelling, low-resource adaptation, evaluation design, robustness analysis, data-centric approaches, and practical deployment challenges. Together, they show that progress in speech technology cannot be measured only by performance on high-resource benchmarks. It must also be assessed by how reliably systems work under realistic constraints, how transparently they are evaluated, and how equitably they serve diverse speakers.

The workshop program includes oral and poster presentations, as well as an invited talk by Jordi Luque from Telefónica Research. We are grateful to our Program Committee, consisting of confirmed domain experts from academia and industry, for their careful and constructive reviews. Their work was essential in shaping a balanced and high-quality program. We also thank all authors for submitting their work, revising their papers, and contributing to the scientific scope of this first edition.

We hope that SPEAKABLE 2026 will serve not only as a venue for presenting current research, but also as a catalyst for collaboration, shared evaluation practices, and community-building around inclusive speech technologies. Our broader aim is captured by the workshop’s guiding message: build strong models, measure what matters, and make bias analysis routine for speech in the long tail.

Finally, we thank the LREC 2026 workshop chairs and organizers for hosting SPEAKABLE as part of the LREC 2026 workshop program. We also thank our invited speaker, reviewers, authors, participants, and supporting institutions for helping make this first edition possible.

Further information about the workshop, including the program and updates, is available on the SPEAKABLE 2026 website: <https://speakable-2026.github.io/>.

The SPEAKABLE 2026 Organizing Committee

Organizing Committee

- Nina Hosseini-Kivanani, Radio Télévision Luxembourg & University of Luxembourg, Luxembourg
- Alessio Brutti, Fondazione Bruno Kessler, Italy
- Marco Matassoni, Fondazione Bruno Kessler, Italy
- Sandipana Dowerah, Tallinn University of Technology, Estonia
- Davide Liga, University of Luxembourg, Luxembourg
- Christoph Schommer, University of Luxembourg, Luxembourg

Program Committee

- Badr M. Abdullah, Saarland University, Germany
- Şeymanur Aktı, Karlsruhe Institute of Technology (KIT), Germany
- Tanel Alumäe, Tallinn University of Technology, Estonia
- Dimitra Anastasiou, Luxembourg Institute of Science and Technology, Luxembourg
- Leonardo Badino, Almagest, Italy
- Stefano Bannò, University of Cambridge, UK
- Carlos Carvalho, INESC-ID, Portugal
- Shammur Absar Chowdhury, Qatar Computing Research Institute (QCRI), Qatar
- Matt Coler, University of Groningen, Netherlands
- Lorenzo Concina, Fondazione Bruno Kessler (FBK), Italy
- Miguel Couceiro, Universidade de Lisboa, Portugal
- Rohan Kumar Das, Fortemedia, Singapore
- Ioannis Douros, Stavros Niarchos Foundation (SNF), Greece
- Yassine El Kheir, Deutsches Forschungszentrum für Künstliche Intelligenz (DFKI), Germany
- Daniele Falavigna, Fondazione Bruno Kessler (FBK), Italy
- Saeed Farzi, Fondazione Bruno Kessler (FBK), Italy
- Maxime Fily, Inalco (Institut national des langues et civilisations orientales), France
- Peter Gilles, University of Luxembourg, Luxembourg
- Nabarun Goswami, University of Tokyo, Japan
- Felix Herron, Université Paris Dauphine – PSL, France
- Aditya Joshi, UNSW Sydney, Australia
- Joonas Kalda, Pyannote.ai, France
- Ajinkya Kulkarni, Idiap Research Institute, Switzerland
- Spyretta Leivaditi, University of Groningen, Netherlands
- Damien Lolive, IUT of Vannes (University of South Brittany), France
- Keerthana Murugaraj, University of Luxembourg, Luxembourg
- Maria Onoeva, Charles University, Czech Republic
- Ludovica Pannitto, University of Bologna, Italy
- Fernando Perez Tellez, Technological University Dublin, Ireland
- Ben Peters, INESC-ID & Instituto Superior Técnico, Portugal

- Fred Philippy, SnT, University of Luxembourg, Luxembourg
- Bornali Phukon, University of Illinois Urbana-Champaign, USA
- Tina Raissi, RWTH Aachen University, Germany
- Thomas Rolland, Orange, France
- Beatrice Savoldi, Fondazione Bruno Kessler (FBK), Italy
- Imran Sheikh, Vivoka, France
- Ravi Shekhar, University of Essex, UK
- Golshid Shekoufandeh, University of Amsterdam, Netherlands
- Francisco Teixeira, INESC-ID, Portugal
- Hawau Olamide Toyin, Mohamed bin Zayed University of Artificial Intelligence (MBZUAI), UAE
- Preben Vangberg, Bangor University, UK
- Jelena Vasić, Technological University Dublin, Ireland
- Martijn Wieling, University of Groningen, Netherlands
- Enrico Zovato, Almagest, Italy
- Juan Pablo Zuluaga-Gomez, AGIGO, Switzerland

Invited speaker

- Jordi Luque, Telefónica, Spain

Table of Contents

<i>Adapting Foundational ASR Models to Efik: An Empirical Study of an Extremely Low-Resource Tonal Language</i>	
Offiong Bassey Edet, Stephen Orok Duke, Enoima Essien Umoh, Benjamin Okon Nyong and Andrew Asuquo Nkpanam	1
<i>PAREDA: A Multi-Accent Speech Dataset of Natural Language Processing Research Discussions</i>	
Sicheng Jin, Dipankar Srirag and Aditya Joshi	8
<i>Say Again? The Limits of Whisper with Conversation. A Case Study on the KIParla Corpus.</i>	
Martina Simonotti, Ludovica Pannitto, Caterina Mauri, Adriano Ferraresi and Gabriele Carlioli	16
<i>Not All Polar Questions Are the Same: ASR, Humans, and Russian</i>	
Maria Onoeva	31
<i>Quantizing Whisper: How Design Choices Affect ASR Performance</i>	
Arthur Söhler, Julian Irigoyen and Andreas Søeborg Kirkedal	39
<i>"OK Aura, Be Fair with Me": Demographics-Agnostic Training for Bias Mitigation in Wake-up Word Detection</i>	
Fernando López, Paula Delgado-Santos, Pablo Gómez, David Solans and Jordi Luque	47
<i>Scalable Expansion of Multilingual Speech LLMs for ASR: A Continual Learning Approach</i>	
Lorenzo Concina, Marco Matassoni and Alessio Brutti	59
<i>Responsible Benchmarking of Fairness for Automatic Speech Recognition</i>	
Felix E. Herron, Ange Richard, François Portet, Alexandre Allauzen and Solange Rossato	66
<i>Addressing Accent Disparities in Automatic Speech Recognition: A Comparative Study of Single and Two-Step Adaptation</i>	
Mykhailo Danilevskyi, Fernando Perez-Tellez and Jelena Vasic	79
<i>Investigating Speaker Pronunciation Variability in Speech Embeddings: Speaker and L1 Effects on French as a Second Language</i>	
Maxime Fily, Martine Adda-Decker and Guillaume Wisniewski	86
<i>What LID Systems Say About Dialectal Variation. The Case of Yiddish, Quechua and Mande</i>	
Johanna Cordova, Eric Jordan and Valentina Fedchenko	98
<i>HARNESS: Lightweight Distilled Arabic Speech Foundation Models</i>	
Vrunda Nileshkumar Sukhadia and Shammur Absar Chowdhury	109
<i>When Does OmniASR Fail? A Fine-Grained Human Evaluation on Saudi Arabic Dialects</i>	
Hend Al-Khalifa	118
<i>SpeechLM for Automatic Speech Recognition in Low-resource Languages</i>	
Md Abdur Razzaq Riyadh, Eneko Agirre, Eva Navas and Claudia Borg	125

<i>Improving Low-resource ASR Using Bilingual Fine-tuning with Language Identification: A Cross-linguistic Evaluation</i>	
Reihaneh Amooie, Yun Hao, Wietse de Vries, Jelske Dijkstra, Matt Coler and Martijn Wieling	132
<i>Leveraging Speech Models for Audio-based Lexical Retrieval in Dictionaries: The Case of the Teochew Language</i>	
Siman Chen, Ilaine Wang, Maxime Fily and Pierre Magistry	139
<i>Stage-Aware Cross-Lingual Transfer for Faroese ASR: When and Which Languages Matter</i>	
Dávid í Lág, Barbara Scalvini, Carlos Daniel Mena and Jón Guðnason	150
<i>Doing More with Less: Determining Optimal Pre-training Model for Irish Automatic Speech Recognition through Multi-step Fine-tuning</i>	
Caoilfhionn Ní Dheoráin, Ruth Holmes, Nicholas Evans, Thomas Laurent, Anthony Ventresque and Ellen Rushe	162
<i>Blank-Aware Decoding for Transcript-Free Phoneme Alignment in Low-Resource Languages and Dialects</i>	
Domenico De Cristofaro, Barbara Plank and Alessandro Vietti	174
<i>On the Role of Encoder Depth: Pruning Whisper and LoRA Fine-Tuning in SLAM-ASR</i>	
Ganesh Pavan Kartikeya Bharadwaj Kolluri, Michael Kampouridis and Ravi Shekhar .	183
<i>TaLK-Corpus: A Regionally Diverse Evaluation Set for Sri Lankan Tamil Speech</i>	
Adsajan Thillainathan, Nishanthini Kanthakumar, Nivethiga Rasan and Kengatharaiyer Sarveswaran	194

Workshop Program

May 11, 2026

Room: W11

09:00–09:20 **Introduction and general remarks**

09:20–10:20 **Invited speaker: Jordi Luque**

10:20–10:30 **Remote posters**

Adapting Foundational ASR Models to Efik: An Empirical Study of an Extremely Low-Resource Tonal Language

Offiong Bassey Edet, Stephen Orok Duke, Enoima Essien Umoh, Benjamin Okon Nyong and Andrew Asuquo Nkpanam

PAREDA: A Multi-Accent Speech Dataset of Natural Language Processing Research Discussions

Sicheng Jin, Dipankar Srirag and Aditya Joshi

10:30–12:00 **Coffee Break and Poster session**

Say Again? The Limits of Whisper with Conversation. A Case Study on the KIParla Corpus.

Martina Simonotti, Ludovica Pannitto, Caterina Mauri, Adriano Ferraresi and Gabriele Carioli

Not All Polar Questions Are the Same: ASR, Humans, and Russian

Maria Onoeva

Quantizing Whisper: How Design Choices Affect ASR Performance

Arthur Söhler, Julian Irigoyen and Andreas Søeborg Kirkedal

"OK Aura, Be Fair with Me": Demographics-Agnostic Training for Bias Mitigation in Wake-up Word Detection

Fernando López, Paula Delgado-Santos, Pablo Gómez, David Solans and Jordi Luque

Scalable Expansion of Multilingual Speech LLMs for ASR: A Continual Learning Approach

Lorenzo Concina, Marco Matassoni and Alessio Brutti

May 11, 2026 (continued)

Responsible Benchmarking of Fairness for Automatic Speech Recognition

Felix E. Herron, Ange Richard, François Portet, Alexandre Allauzen and Solange Rossato

Addressing Accent Disparities in Automatic Speech Recognition: A Comparative Study of Single and Two-Step Adaptation

Mykhailo Danilevskyi, Fernando Perez-Tellez and Jelena Vasic

Investigating Speaker Pronunciation Variability in Speech Embeddings: Speaker and L1 Effects on French as a Second Language

Maxime Fily, Martine Adda-Decker and Guillaume Wisniewski

What LID Systems Say About Dialectal Variation. The Case of Yiddish, Quechua and Mande

Johanna Cordova, Eric Jordan and Valentina Fedchenko

HARNESS: Lightweight Distilled Arabic Speech Foundation Models

Vrunda Nileshekumar Sukhadia and Shammur Absar Chowdhury

When Does OmniASR Fail? A Fine-Grained Human Evaluation on Saudi Arabic Dialects

Hend Al-Khalifa

12:00–13:00

Spotlight papers

SpeechLM for Automatic Speech Recognition in Low-resource Languages

Md Abdur Razzaq Riyadh, Eneko Agirre, Eva Navas and Claudia Borg

Improving Low-resource ASR Using Bilingual Fine-tuning with Language Identification: A Cross-linguistic Evaluation

Reihaneh Amooie, Yun Hao, Wietse de Vries, Jelske Dijkstra, Matt Coler and Martijn Wieling

May 11, 2026 (continued)

13:00–14:00 Lunch break

14:00–16:00 Architectures and learning methods

Leveraging Speech Models for Audio-based Lexical Retrieval in Dictionaries: The Case of the Teochew Language

Siman Chen, Ilaine Wang, Maxime Fily and Pierre Magistry

Stage-Aware Cross-Lingual Transfer for Faroese ASR: When and Which Languages Matter

Dávid í Lág, Barbara Scavini, Carlos Daniel Mena and Jón Guðnason

Doing More with Less: Determining Optimal Pre-training Model for Irish Automatic Speech Recognition through Multi-step Fine-tuning

Caoilfhionn Ní Dheoráin, Ruth Holmes, Nicholas Evans, Thomas Laurent, Anthony Ventresque and Ellen Rushe

Blank-Aware Decoding for Transcript-Free Phoneme Alignment in Low-Resource Languages and Dialects

Domenico De Cristofaro, Barbara Plank and Alessandro Vietti

On the Role of Encoder Depth: Pruning Whisper and LoRA Fine-Tuning in SLAM-ASR

Ganesh Pavan Kartikeya Bharadwaj Kolluri, Michael Kampouridis and Ravi Shekhar

TaLK-Corpus: A Regionally Diverse Evaluation Set for Sri Lankan Tamil Speech

Adsajan Thillainathan, Nishanthini Kanthakumar, Nivethiga Rasan and Kengatharaiyer Sarveswaran

May 11, 2026 (continued)

16:00–16:30 Coffee break

16:30–17:00 Best paper and closing remarks