

Table of Contents

Error Analysis in a Modular Meeting Transcription System	3
<i>Peter Vieting, Simon Berger, Thilo von Neumann, Christoph Boeddeker, Ralf Schlüter, Reinhold Haeb-Umbach</i>	
Building a German-centric SpeechLLM Using Limited Data	8
<i>Manas Maurya, Thomas Dethmann, Oliver Walter, Christoph Andreas Schmidt, Joachim Köhler</i>	
Investigation of Speech and Noise Latent Representations in Single-channel VAE-based Speech Enhancement	13
<i>Jiatong Li, Simon Doclo</i>	
Unveiling Deep Speech Embeddings: Acoustic Insights into Hatespeech Detection	18
<i>Rathi Adarshi Rammohan, Zhao Ren, Aleksandra Świdarska, Dennis Küster, Tanja Schultz</i>	
The Role of Surprisal and Entropy in Spoken Intercomprehension: An Experiment on Translation of Cognates with Varied Predictability	23
<i>Wei Xue, Julius Steuer, Dietrich Klakow, Bernd Möbius</i>	
Extending Manifold-Based MIMO System Identification to Adaptive Crosstalk Cancellation	28
<i>Johannes Hahn, Tobias Kabzinski, Peter Jax</i>	
Speaker vs Noise Conditioning for Adaptive Speech Enhancement	33
<i>Andreas Triantafyllopoulos, Iosif Tsangko, Michael Müller, Hendrik Schröter, Björn Schuller</i>	
YOLO-based Signal Detection of Amplitude Modulated Audio Transmissions in Realistic HF Scenarios	38
<i>Lukas Henneke, Sebastian Urrigshardt, Fabian Fritz, Frank Kurth</i>	
Effectiveness of Acceleration Sensors on the Thorax and Abdomen for Speech Breathing Analysis	43
<i>Dani Kazzy, Christian Kleiner, Susanne Fuchs, Peter Birkholz</i>	
Exploring In-Context Learning Capabilities of ChatGPT for Pathological Speech Detection	48
<i>Mahdi Amiri, Hatef Otroushi Shahreza, Ina Kodrasi</i>	
Navigating PESQ: Up-to-Date Versions and Open Implementations	53
<i>Matteo Torcoli, Mhd Modar Halimeh, Emanuel A. P. Habets</i>	
Mixed-Effects Models Neural Networks for Improved Speech-Based Predictions of Cognitive Decline	58
<i>Jordan Behrendt, jiumend Zhang, Claudia Börnhorst, Tanja Schultz</i>	
Personalized Speech Synthesis for Zero-Shot Keyword Spotting	63
<i>Fahrettin Gökgöz, Alessia Cornaggia-Urrigshardt, Kevin Wilkinghoff</i>	
Target Speaker Extraction: the Importance of a Powerful Extractor and Content-Informed Embeddings	68
<i>Elias De Souter, Stijn Kindt, Kaixuan Yang, Haixin Zhao, Siyuan Song, Yanjue Song, Nilesch Madhu</i>	
Enhancement of Neural Embeddings for Speaker Identification in Ad-hoc Acoustic Sensor Networks and Multi-Speaker Scenarios	73
<i>Philipp Intek, Luca Becker, Timm Koppelman, Rainer Martin</i>	
A fully Zero-shot Approach to Obtaining Specialized and Compact Audio Tagging Models	78
<i>Alexander Werning, Reinhold Häb-Umbach</i>	
Evaluating the Impact of Crowdsourced Audio Data on Speech Quality Assessment	83
<i>Kirill Shchegelskiy, Malek El-Tannir, Wafaa Wardah, Tuğçe Melike Koçak Büyüktas, Sebastian Möller</i>	
A Comparative Analysis on ASR System Combination for Attention, CTC, Factored Hybrid, and Transducer Models	88
<i>Noureddin Bayoumi, Robin Schmitt, Tina Raissi, Albert Zeyer, Ralf Schlüter, Hermann Ney</i>	
Neural Prosody Prediction for German Articulatory Speech Synthesis	93
<i>Peter Steiner, Zihao Huang, Arne-Lukas Fietkau, Peter Birkholz</i>	

Adapting the Fréchet Audio Distance as an Objective Metric for Text-to-Speech Quality Evaluation	98
<i>Laurymas Zavistanavicius, Frank Zalkow, Christian Dittmar, Robert L. Stevenson</i>	
Towards Complex-Valued VAE-Based Distillation for Representation Learning in Speech Enhancement	103
<i>Haixin Zhao, Kaixuan Yang, Nilesch Madhu</i>	
Room Reverberation Effectively Masks Deepfake Traces	108
<i>Sophie Hoppe, Anabell Hacker, Markus Brückl</i>	
Early and Late Reflections in Acoustic Echo Control: An Experimental Study on (Neural) Kalman Filters and DNN Methods	113
<i>Ernst Seidel, Tim Fingscheidt</i>	
On the Application of Diffusion Models for Simultaneous Denoising and Dereverberation	118
<i>Adrian Meise, Tobias Cord-Landwehr, Reinhold Haeb-Umbach</i>	
Binaural Distance Estimation Using a Joint Latent Representation of Acoustic Distance and Direct Path Response	123
<i>Daniel Neudek, Benjamin Stodt, Stephan Getzmann, Rainer Martin</i>	
Detecting COPD Exacerbations Before Onset Using Vocal Biomarkers	128
<i>Lars Nippert, Sami O. Simons, Julia Hoxha</i>	
Comparison of Knowledge Distillation Methods for Low-complexity Multi- microphone Speech Enhancement using the FT-JNF Architecture	133
<i>Robert Metzger, Mattes Ohlenbusch, Christian Rollwage, Simon Doclo</i>	
Acoustical characterization and perceptual comparison of four types of 3D- printed vocal tract models for the German and Japanese vowels /a,e,i,o,u/	138
<i>Christian Kleiner, Peter Birkholz, Dominik Schäfer, Takayuki Arai</i>	
Low-Complexity Neural Wind Noise Reduction for Audio Recordings	143
<i>Hesam Eftekhari, Srikanth Raj Chetupalli, Shrishti Saha Shetu, Emanuël A. P. Habets, Oliver Thiergart</i>	
An Improved Neural Network Architecture for Target Speech Extraction	148
<i>David Joos, Friedrich Faubel, Jonas Jungclaussen, Markus Buck, Wolfgang Minker</i>	
A Very-Low Delay High-Performance Speech Vocoder Based on the Encodec Speech Decoder	153
<i>Renzheng Shi, Tim Fingscheidt</i>	
Evaluating the Recognition Performance of the RehaLingo Speech Training System with Aphasic Speech	158
<i>Hans-Günther Hirsch, Yannic Tiggelkamp, Christian Neumann, Tobias Bolten</i>	
Optimization of Feature and Loss Exponents for Lightweight DNN-based Binaural Speech Enhancement	163
<i>Aleksej Chinaev, Gerald Enzner, Stefan Thaleiser</i>	
Unified Learnable 2D Convolutional Feature Extraction for ASR	168
<i>Peter Vieting, Benedikt Hilmes, Ralf Schlüter, Hermann Ney</i>	
Blind Estimation of Head Rotations From Binaural Recordings	173
<i>Erik Fleischhauer, Peter Jax</i>	
ReverbFX: A Dataset of Room Impulse Responses Derived from Reverb Effect Plugins for Singing Voice Dereverberation	178
<i>Julius Richter, Till Svajda, Timo Gerkmann</i>	