

ICASSP 2022 - 2022 IEEE International Conference on Acoustics, Speech and Signal Processing

**Virtual Conference
7-13 May 2022**

**Singapore
22-27 May 2022**

Pages 1-535



**IEEE Catalog Number: CFP22ICA-POD
ISBN: 978-1-6654-0541-6**

**Copyright © 2022 by the Institute of Electrical and Electronics Engineers, Inc.
All Rights Reserved**

Copyright and Reprint Permissions: Abstracting is permitted with credit to the source. Libraries are permitted to photocopy beyond the limit of U.S. copyright law for private use of patrons those articles in this volume that carry a code at the bottom of the first page, provided the per-copy fee indicated in the code is paid through Copyright Clearance Center, 222 Rosewood Drive, Danvers, MA 01923.

For other copying, reprint or republication permission, write to IEEE Copyrights Manager, IEEE Service Center, 445 Hoes Lane, Piscataway, NJ 08854. All rights reserved.

****** This is a print representation of what appears in the IEEE Digital Library. Some format issues inherent in the e-media version may also appear in this print version.***

IEEE Catalog Number:	CFP22ICA-POD
ISBN (Print-On-Demand):	978-1-6654-0541-6
ISBN (Online):	978-1-6654-0540-9
ISSN:	1520-6149

Additional Copies of This Publication Are Available From:

Curran Associates, Inc
57 Morehouse Lane
Red Hook, NY 12571 USA
Phone: (845) 758-0400
Fax: (845) 758-2633
E-mail: curran@proceedings.com
Web: www.proceedings.com

CURRAN ASSOCIATES INC.
proceedings
.com

TABLE OF CONTENTS

ASPS-1: SIGNAL PROCESSING FOR HEALTH AND WELLNESS MONITORING

ASPS-1.1: COUGHTRIGGER: EARBUDS IMU BASED COUGH DETECTION 1 ACTIVATOR USING AN ENERGY-EFFICIENT SENSITIVITY-PRIORITIZED TIME SERIES CLASSIFIER <i>Shibo Zhang, Nabil Alshurafa, Northwestern University, United States of America; Ebrahim Nemati, Tousif Ahmed, Mahbubur Rahman, Jilong Kuang, Alex Gao, Samsung Research America, United States of America; Minh Dinh, Nathan Folkman, Samsung Design Innovation Center, United States of America</i>	
ASPS-1.2: NON-INVASIVE BLOOD PRESSURE MONITORING WITH MULTI-MODAL 6 IN-EAR SENSING <i>Hoang Truong, Alessandro Montanari, Fahim Kawsar, Nokia Bell Labs, United Kingdom of Great Britain and Northern Ireland</i>	
ASPS-1.3: INTELLIGENT WI-FI BASED CHILD PRESENCE DETECTION SYSTEM 11 <i>Xiaolu Zeng, K. J. Ray Liu, University of Maryland College Park, United States of America; Beibei Wang, Sai Deepika Regani, Origin Wireless Inc, United States of America; Chenshu Wu, The University of Hong Kong, Hong Kong</i>	
ASPS-1.4: REAL-TIME FALL DETECTION USING MMWAVE RADAR 16 <i>Wenxuan Li, Dongheng Zhang, Yadong Li, Zhi Wu, Jinbo Chen, Dong Zhang, Yang Hu, Qibin Sun, Yan Chen, University of Science and Technology of China, China</i>	
ASPS-1.5: HIERARCHICAL DEEP LEARNING MODEL WITH INERTIAL AND 21 PHYSIOLOGICAL SENSORS FUSION FOR WEARABLE-BASED HUMAN ACTIVITY RECOGNITION <i>Dae Yon Hwang, Pai Chet Ng, Dimitrios Hatzinakos, Konstantinos N. Plataniotis, University of Toronto, Canada; Yuanhao Yu, Yang Wang, Huawei Technologies, Canada; Petros Spachos, University of Guelph, Canada</i>	

ASPS-2: AUDIO AND ACOUSTIC SIGNAL PROCESSING

ASPS-2.1: SPEECH RECOVERY FOR REAL-WORLD SELF-POWERED 26 INTERMITTENT DEVICES <i>Yu-Chen Lin, Tsun-An Hsieh, Kuo-Hsuan Hung, Cheng Yu, Yu Tsao, Academia Sinica, Taiwan; Harinath Garudadri, University of California, San Diego, United States of America; Tei-Wei Kuo, City University of Hong Kong, Hong Kong</i>	
ASPS-2.2: PHASE CONTROL OF PARAMETRIC ARRAY LOUDSPEAKER BY 31 OPTIMIZING SIDEBAND WEIGHTS <i>Ai Okano, Yoshinobu Kajikawa, Kansai University, Japan</i>	
ASPS-2.3: LOW-LATENCY HUMAN-COMPUTER AUDITORY INTERFACE BASED ON 36 REAL-TIME VISION ANALYSIS <i>Florian Scalvini, Cyrille Migniot, Julien Dubois, ImViA, France; Camille Bordeau, Maxime Ambard, LEAD, France</i>	
ASPS-2.4: ROBUST ADAPTIVE NOISE CANCELLER ALGORITHM WITH SNR-BASED 41 STEPSIZE CONTROL AND NOISE-PATH GAIN COMPENSATION <i>Akihiko Sugiyama, Yahoo Japan Corporation, Japan</i>	
ASPS-2.5: NEARTRACKER: ACOUSTIC 2-D TARGET TRACKING WITH NEARBY 46 REFLECTOR IN SISO SYSTEM <i>Chao Liu, Linlin Gao, Ruobing Jiang, Ocean University of China, China</i>	

ASPS-3: SIGNAL PROCESSING, MACHINE LEARNING, AND ARCHITECTURE DESIGN OPTIMIZATION

ASPS-3.1: AN EFFICIENT METHOD FOR GENERIC DSP IMPLEMENTATION OF DILATED CONVOLUTION 51

Harinarayanan EV, Sachin Ghanekar, Cadence Design Systems, India

ASPS-3.2: COMPRESSION-AWARE PROJECTION WITH GREEDY DIMENSION REDUCTION FOR CONVOLUTIONAL NEURAL NETWORK ACTIVATIONS 56

Yu-Shan Tai, Chieh-Fang Teng, Cheng-Yang Chang, An-Yeu Wu, National Taiwan University, Taiwan

ASPS-3.3: OPTIMIZING THE CONSUMPTION OF SPIKING NEURAL NETWORKS WITH ACTIVITY REGULARIZATION 61

Simon Narduzzi, Siavash A. Bigdeli, L. Andrea Dunbar, Centre suisse d'électronique et de microtechnique, Switzerland; Shih-Chii Liu, University of Zürich, Switzerland

ASPS-3.4: IMPQ: REDUCED COMPLEXITY NEURAL NETWORKS VIA GRANULAR PRECISION ASSIGNMENT 66

Sujan Kumar Gonugondla, Amazon, United States of America; Naresh Shanbhag, University of Illinois at Urbana-Champaign, United States of America

ASPS-3.5: RATE CODING OR DIRECT CODING: WHICH ONE IS BETTER FOR ACCURATE, ROBUST, AND ENERGY-EFFICIENT SPIKING NEURAL NETWORKS? 71

Youngeun Kim, Hyoungeob Park, Abhishek Moitra, Abhiroop Bhattacharjee, Yeshwanth Venkatesha, Priyadarshini Panda, Yale University, United States of America

ASPS-3.6: PYXIS: AN OPEN-SOURCE PERFORMANCE DATASET OF SPARSE ACCELERATORS 76

Linghao Song, Yuze Chi, Jason Cong, University of California, Los Angeles, United States of America

ASPS-4: INTEGRATIVE SIGNAL PROCESSING AND MACHINE LEARNING FOR MULTIMODAL SENSING AND PROCESSING

ASPS-4.1: FAST FAULT DIAGNOSIS METHOD OF ROLLING BEARINGS IN MULTI-SENSOR MEASUREMENT ENVIRONMENT 81

Zuozhou Pan, Zong Meng, Yanshan University, China; Zhiping Lin, Yuanjin Zheng, Nanyang Technological University, Singapore

ASPS-4.2: DETECTING ANOMALY IN CHEMICAL SENSORS VIA REGULARIZED CONTRASTIVE LEARNING 86

Diaa Badawi, Ahmet Enis Cetin, University of Illinois at Chicago, United States of America; Ishaan Bassi, Sule Ozev, Arizona State University, United States of America

ASPS-4.3: EVOLUTIONARY NEURAL ARCHITECTURE DESIGN OF LIQUID STATE MACHINE FOR IMAGE CLASSIFICATION 91

Cheng Tang, University of Toyama, Japan; Junkai Ji, Qiuzhen Lin, Shenzhen University, China; Yan Zhou, Northeastern University, China

ASPS-4.4: INVISIBLE AND EFFICIENT BACKDOOR ATTACKS FOR COMPRESSED DEEP NEURAL NETWORKS 96

Huy Phan, Yi Xie, Yingying Chen, Bo Yuan, Rutgers University, United States of America; Jian Liu, University of Tennessee, United States of America

ASPS-4.5: TENSOR-BASED ORTHOGONAL MATCHING PURSUIT WITH PHASE ROTATION FOR CHANNEL ESTIMATION IN HYBRID BEAMFORMING MIMO-OFDM SYSTEMS 101

Cheng-Hung Lo, Pei-Yun Tsai, National Central University, Taiwan, Province of China

AUD-1: MUSIC SOURCE SEPARATION I: INTERACTIVE, CONDITIONED, AND INFORMED MODELS

AUD-1.1: SPAIN-NET: SPATIALLY-INFORMED STEREOPHONIC MUSIC SOURCE SEPARATION 106

Darius Petermann, Minje Kim, Indiana University, United States of America

AUD-1.3: IMPROVED SINGING VOICE SEPARATION WITH CHROMAGRAM-BASED PITCH-AWARE REMIXING 111

Siyuan Yuan, Stanford University, United States of America; Zhepei Wang, UIUC, United States of America; Umut Isik, Ritwik Giri, Jean-Marc Valin, Michael M. Goodwin, Arvinth Krishnaswamy, Amazon Web Services, United States of America

AUD-1.4: DON'T SEPARATE, LEARN TO REMIX: END-TO-END NEURAL REMIXING WITH JOINT OPTIMIZATION 116

Haici Yang, Shivani Firodiya, Minje Kim, Indiana University, United States of America; Nicholas Bryan, Adobe Research, United States of America

AUD-1.5: FEW-SHOT MUSICAL SOURCE SEPARATION 121

Yu Wang, Juan Pablo Bello, New York University, United States of America; Daniel Stoller, Rachel M. Bittner, Spotify, United States of America

AUD-1.6: SOURCE SEPARATION BY STEERING PRETRAINED MUSIC MODELS 126

Ethan Manilow, Patrick O'Reilly, Bryan Pardo, Northwestern University, United States of America; Prem Seetharaman, Descript, Inc., United States of America

AUD-2: DETECTION AND CLASSIFICATION OF ACOUSTIC SCENES AND EVENTS I: DATASETS

AUD-2.1: INFANT CRYING DETECTION IN REAL-WORLD ENVIRONMENTS..... 131

Xuwen Yao, Megan Micheletti, Mckensey Johnson, Edison Thomaz, Kaya de Barbaro, University of Texas at Austin, United States of America

AUD-2.2: WIKITAG: WIKIPEDIA-BASED KNOWLEDGE EMBEDDINGS TOWARDS IMPROVED ACOUSTIC EVENT CLASSIFICATION 136

Qin Zhang, Qingming Tang, Chieh-Chi Kao, Ming Sun, Yang Liu, Chao Wang, Amazon Inc, United States of America

AUD-2.3: URBAN SOUND & SIGHT: DATASET AND BENCHMARK FOR AUDIO-VISUAL URBAN SCENE UNDERSTANDING 141

Magdalena Fuentes, Bea Steers, Julia Wilkins, Qianyi Shi, Yao Hou, Juan Pablo Bello, New York University, United States of America; Pablo Zinemanas, Xavier Serra, Universitat Pompeu Fabra, Spain; Martin Rocamora, Universidad de la Republica, Uruguay; Luca Bondi, Samarjit Das, Bosch Research, United States of America

AUD-2.4: REAL-WORLD ON-BOARD UAV AUDIO DATA SET FOR PROPELLER ANOMALIES 146

Sai Srinadhu Katta, Kide Vuojärvi, Sivaprasad Nandyala, Ulla-Maria Kovalainen, Lauren Baddeley, Secure Systems Research Centre (SSRC) - Technology Innovation Institute (TII), United Arab Emirates

AUD-2.5: VOCALSOUND: A DATASET FOR IMPROVING HUMAN VOCAL SOUNDS RECOGNITION 151

Yuan Gong, James Glass, Massachusetts Institute of Technology, United States of America; Jin Yu, Signify Research, United States of America

AUD-2.6: WEARABLE SELD DATASET: DATASET FOR SOUND EVENT LOCALIZATION AND DETECTION USING WEARABLE DEVICES AROUND HEAD 156

Kento Nagatomo, Kohei Yatabe, Yasuhiro Oikawa, Waseda University, Japan; Masahiro Yasuda, Shoichiro Saito, NTT Corporation, Japan

AUD-3: SYSTEM IDENTIFICATION AND REVERBERATION REDUCTION

AUD-3.1: TUNET: A BLOCK-ONLINE BANDWIDTH EXTENSION MODEL BASED ON TRANSFORMERS AND SELF-SUPERVISED PRETRAINING 161

Viet-Anh Nguyen, Anh H. T. Nguyen, FPT Software, Viet Nam; Andy W. H. Khong, Nanyang Technological University, Singapore

AUD-3.2: DRC-NET: DENSELY CONNECTED RECURRENT CONVOLUTIONAL NEURAL NETWORK FOR SPEECH DEREVERBERATION 166

Jinjiang Liu, Xueliang Zhang, College of Computer Science, Inner Mongolia University, China

AUD-3.3: CUSTOMIZABLE END-TO-END OPTIMIZATION OF ONLINE NEURAL NETWORK-SUPPORTED DEREVERBERATION FOR HEARING DEVICES 171

Jean-Marie Lemerrier, Timo Gerkmann, Universität Hamburg, Germany; Joachim Thiemann, Raphael Koning, Advanced Bionics GmbH, Germany

AUD-3.5: IMPORTANCE OF SWITCH OPTIMIZATION CRITERION IN SWITCHING WPE DEREVERBERATION 176

Naoyuki Kamo, Rintaro Ikeshita, Keisuke Kinoshita, Tomohiro Nakatani, NTT, Japan

AUD-4: SYMBOLIC ANALYSIS AND SYNTHESIS OF MUSIC

AUD-4.1: AUDIO-TO-SYMBOLIC ARRANGEMENT VIA CROSS-MODAL MUSIC REPRESENTATION LEARNING 181

Ziyu Wang, Gus Xia, NYU Shanghai, China; Dejing Xu, Tencent ARC Lab, China; Ying Shan, Tencnet ARC Lab, China

AUD-4.2: MUSIC PHRASE INPAINTING USING LONG-TERM REPRESENTATION AND CONTRASTIVE LOSS 186

Shiqi Wei, Weiguo Gao, Fudan University, China; Gus Xia, Liwei Lin, New York University Shanghai, China; Yixiao Zhang, Queen Mary University of London, United Kingdom of Great Britain and Northern Ireland

AUD-4.3: MELONS: GENERATING MELODY WITH LONG-TERM STRUCTURE USING TRANSFORMERS AND STRUCTURE GRAPH 191

Yi Zou, Communication University of China, China; Pei Zou, Peking University, China; Yi Zhao, Ran Zhang, Xiaorui Wang, Kuaishou Technology Co., Beijing, China; Kaixiang Zhang, Kuaishou Technology Co, China

AUD-4.4: DIFFICULTY-AWARE NEURAL BAND-TO-PIANO SCORE ARRANGEMENT BASED ON NOTE- AND STATISTIC-LEVEL CRITERIA 196

Moyu Terao, Yuki Hiramatsu, Ryoto Ishizuka, Yiming Wu, Kazuyoshi Yoshii, Kyoto University, Japan

AUD-4.5: SCORE DIFFICULTY ANALYSIS FOR PIANO PERFORMANCE EDUCATION BASED ON FINGERING 201

Pedro Ramoneda, Music Technology Group (MTG), Spain; Nazif Can Tamer, Vsevolod Eremenko, Xavier Serra, Marius Mirón, MTG, Spain

AUD-5: DETECTION AND CLASSIFICATION OF ACOUSTIC SCENES AND EVENTS II: TASKS

AUD-5.1: A NEURAL NETWORK-BASED HOWLING DETECTION METHOD FOR REAL-TIME COMMUNICATION APPLICATIONS 206

Zhipeng Chen, Yiya Hao, Yaobin Chen, Gong Chen, Liang Ruan, Netease GrowthEase, China

AUD-5.2: ALARM SOUND DETECTION USING TOPOLOGICAL SIGNAL PROCESSING 211

Tomer Fireaizen, Saar Ron, Omer Bobrowski, Technion - Israel Institute of Technology, Israel

AUD-5.3: A METHOD FOR ESTIMATING THE GROUPING OF PARTICIPANTS IN CLASSROOM GROUP WORK USING ONLY AUDIO INFORMATION 216

Osamu Ichikawa, Yuuto Shima, Shiga University, Japan; Takahiro Nakayama, The University of Tokyo, Japan; Hajime Shirouzu, National Institute for Educational Policy Research, Japan

AUD-5.4: ENVIRONMENTAL SOUND EXTRACTION USING ONOMATOPOEIC WORDS	221
<i>Yuki Okamoto, Ritsumeikan University, Japan; Shota Horiguchi, Masaaki Yamamoto, Yohei Kawaguchi, Hitachi Ltd., Japan; Keisuke Imoto, Doshisha University, Japan</i>	
AUD-5.6: ECHO-AWARE ADAPTATION OF SOUND EVENT LOCALIZATION AND DETECTION IN UNKNOWN ENVIRONMENTS	226
<i>Masahiro Yasuda, Yasunori Ohishi, Shoichiro Saito, NTT Corporation, Japan</i>	
AUD-6: AUDIO SECURITY AND MULTI-MODAL SYSTEMS	
AUD-6.1: ON ADVERSARIAL ROBUSTNESS OF LARGE-SCALE AUDIO VISUAL LEARNING	231
<i>Juncheng B Li, Xinjian Li, Po-Yao (Bernie) Huang, Florian Metze, Carnegie Mellon University, United States of America; Shuhui Qu, Stanford University, United States of America</i>	
AUD-6.2: ADVERSARIAL SAMPLE DETECTION FOR SPEAKER VERIFICATION BY NEURAL VOCODERS	236
<i>Haibin Wu, Po-chun Hsu, Hung-yi Lee, National Taiwan University, China; Ji Gao, Shanshan Zhang, Shen Huang, Jian Kang, Tencent, China; Zhiyong Wu, Shenzhen International Graduate School, Tsinghua University, China; Helen Meng, The Chinese University of Hong Kong, Hong Kong</i>	
AUD-6.3: SOURCE MIXING AND SEPARATION ROBUST AUDIO STEGANOGRAPHY	241
<i>Naoya Takahashi, Mayank Kumar Singh, Yuki Mitsufuji, Sony Group Corporation, Japan</i>	
AUD-6.4: MULTI-MODAL PRE-TRAINING FOR AUTOMATED SPEECH RECOGNITION	246
<i>David Chan, University of California, Berkeley, United States of America; Shalini Ghosh, Debmalaya Chakrabarty, Björn Hoffmeister, Amazon, United States of America</i>	
AUD-6.5: SPEAKER-TARGETED AUDIO-VISUAL SPEECH RECOGNITION USING A HYBRID CTC/ATTENTION MODEL WITH INTERFERENCE LOSS	251
<i>Ryota Tsunoda, Ryoichi Takashima, Tetsuya Takiguchi, Kobe University, Japan; Ryo Aihara, Yoshie Imai, Mitsubishi Electric Corporation, Japan</i>	
AUD-6.6: TIME-DOMAIN AUDIO-VISUAL SPEECH SEPARATION ON LOW QUALITY VIDEOS	256
<i>Yifei Wu, Chenda Li, Yanmin Qian, Shanghai Jiao Tong University, China; Jinfeng Bai, Zhongqin Wu, TAL Education Group, China, China</i>	
AUD-7: SIGNAL ENHANCEMENT AND RESTORATION I: MULTICHANNEL	
AUD-7.1: COMPLEX-VALUED SPATIAL AUTOENCODERS FOR MULTICHANNEL SPEECH ENHANCEMENT	261
<i>Mhd Modar Halimeh, Walter Kellermann, Friedrich–Alexander University Erlangen–Nürnberg, Germany</i>	
AUD-7.2: MULTICHANNEL NOISE REDUCTION USING DILATED MULTICHANNEL U-NET AND PRE-TRAINED SINGLE-CHANNEL NETWORK	266
<i>Zhi-Wei Tan, Yuan Liu, Andy W. H. Khong, Nanyang Technological University, Singapore; Anh H. T. Nguyen, FPT Software, Viet Nam</i>	
AUD-7.3: ONE MODEL TO ENHANCE THEM ALL: ARRAY GEOMETRY AGNOSTIC MULTI-CHANNEL PERSONALIZED SPEECH ENHANCEMENT	271
<i>Hassan Taherian, The Ohio State University, United States of America; Sefik Emre Eskimez, Takuya Yoshioka, Huaming Wang, Zhuo Chen, Xuedong Huang, Microsoft, United States of America</i>	
AUD-7.4: MULTI-CHANNEL SPEECH DENOISING FOR MACHINE EARS	276
<i>Cong Han, Columbia University, United States of America; Merve Kaya, Kyle Hoefer, Simon Carlile, X Development, LLC, United States of America; Malcolm Slaney, Google, LLC, United States of America</i>	

AUD-7.5: LOCALIZATION BASED SEQUENTIAL GROUPING FOR CONTINUOUS SPEECH SEPARATION	281
<i>Zhong-Qiu Wang, DeLiang Wang, The Ohio State University, United States of America</i>	
AUD-7.6: CONVOLUTIONAL WEIGHTED MINIMUM MEAN SQUARE ERROR FILTER FOR JOINT SOURCE SEPARATION AND DEREVERBERATION	286
<i>Mieszko Fraś, Marcin Witkowski, Konrad Kowalczyk, AGH University of Science and Technology, Poland</i>	
AUD-8: MUSIC SOURCE SEPARATION II: METHODS AND ARCHITECTURES	
AUD-8.1: IMPROVING SOURCE SEPARATION BY EXPLICITLY MODELING DEPENDENCIES BETWEEN SOURCES	291
<i>Ethan Manilow, Bryan Pardo, Northwestern University, United States of America; Curtis Hawthorne, Cheng-Zhi Anna Huang, Jesse Engel, Google, United States of America</i>	
AUD-8.2: MUSIC SOURCE SEPARATION WITH DEEP EQUILIBRIUM MODELS	296
<i>Yuichiro Koyama, Naoki Murata, Shusuke Takahashi, Yuki Mitsufuji, Sony Group Corporation, Japan; Stefan Uhlich, Giorgio Fabbro, Sony Europe B.V., Germany</i>	
AUD-8.5: HARMONIC AND PERCUSSIVE SOUND SEPARATION BASED ON MIXED PARTIAL DERIVATIVE OF PHASE SPECTROGRAM	301
<i>Natsuki Akaishi, Kohei Yatabe, Yasuhiro Oikawa, Waseda University, Japan</i>	
AUD-8.6: ON LOSS FUNCTIONS AND EVALUATION METRICS FOR MUSIC SOURCE SEPARATION	306
<i>Enric Gusó, Universitat Pompeu Fabra, Spain; Jordi Pons, Santiago Pascual, Joan Serrà, Dolby Laboratories, Spain</i>	
AUD-9: DETECTION AND CLASSIFICATION OF ACOUSTIC SCENES AND EVENTS III: LOSSES AND TRAINING	
AUD-9.1: TIME-BALANCED FOCAL LOSS FOR AUDIO EVENT DETECTION	311
<i>Sangwook Park, Mounya Elhilali, Johns Hopkins University, United States of America</i>	
AUD-9.2: MULTI-ACCDOA: LOCALIZING AND DETECTING OVERLAPPING SOUNDS FROM THE SAME CLASS WITH AUXILIARY DUPLICATING PERMUTATION INVARIANT TRAINING	316
<i>Kazuki Shimada, Yuichiro Koyama, Shusuke Takahashi, Naoya Takahashi, Emiru Tsunoo, Yuki Mitsufuji, Sony Group Corporation, Japan</i>	
AUD-9.3: IMPROVED REPRESENTATION LEARNING FOR ACOUSTIC EVENT CLASSIFICATION USING TREE-STRUCTURED ONTOLOGY	321
<i>Arman Zharmagambetov, University of California, Merced, United States of America; Qingming Tang, Chieh-Chi Kao, Qin Zhang, Ming Sun, Viktor Rozgic, Jasha Droppo, Chao Wang, Amazon, Alexa, United States of America</i>	
AUD-9.4: TEMPORAL CONTRASTIVE-LOSS FOR AUDIO EVENT DETECTION	326
<i>Sandeep Kothinti, Mounya Elhilali, Johns Hopkins University, United States of America</i>	
AUD-9.5: A FRAME LOSS OF MULTIPLE INSTANCE LEARNING FOR WEAKLY SUPERVISED SOUND EVENT DETECTION	331
<i>Xu Wang, Xiangjinzi Zhang, Shengwu Xiong, School of Computer and Artificial Intelligence, Wuhan University of Technology, Wuhan, China, China; Yunfei Zi, ,</i>	
AUD-9.6: PSEUDO STRONG LABELS FOR LARGE SCALE WEAKLY SUPERVISED AUDIO TAGGING	336
<i>Heinrich Dinkel, Zhiyong Yan, Yongqing Wang, Junbo Zhang, Yujun Wang, Xiaomi Corporation, China</i>	

AUD-10: HEARING INSTRUMENTS AND ON-DEVICE PROCESSING

AUD-10.1: INDIVIDUALIZED HEAR-THROUGH FOR ACOUSTIC TRANSPARENCY 341 USING PCA-BASED SOUND PRESSURE ESTIMATION AT THE EARDRUM

Wenyu Jin, Tim Schoof, Henning Schepker, Starkey Hearing Technologies, United States of America

AUD-10.2: ON SPECTRAL AND TEMPORAL SPARSIFICATION OF SPEECH SIGNALS 346 FOR THE IMPROVEMENT OF SPEECH PERCEPTION IN CI LISTENERS

Benjamin Lentz, Rainer Martin, Ruhr-University Bochum, Germany; Kirsten Oberländer, Christiane Völter, St. Elisabeth-Hospital, Ruhr-Universität Bochum, Germany

AUD-10.3: A DIFFERENTIABLE OPTIMISATION FRAMEWORK FOR THE DESIGN OF 351 INDIVIDUALISED DNN-BASED HEARING-AID STRATEGIES

Fotios Drakopoulos, Sarah Verhulst, Ghent University, Belgium

AUD-10.4: PERSONALIZED SPEECH ENHANCEMENT: NEW MODELS AND 356 COMPREHENSIVE EVALUATION

Sefik Emre Eskimez, Takuya Yoshioka, Huaming Wang, Xiaofei Wang, Zhuo Chen, Xuedong Huang, Microsoft, United States of America

AUD-10.5: DYNAMIC SLIDING WINDOW FOR REALTIME DENOISING NETWORKS 361

Jinxu Xiang, Yuyang Zhu, Rundi Wu, Ruilin Xu, Changxi Zheng, Columbia University, United States of America; Yuko Ishiwaka, SoftBank Group, Japan

AUD-11: DEEP LEARNING-BASED SINGLE-CHANNEL SPEECH ENHANCEMENT

AUD-11.1: BLOOM-NET: BLOCKWISE OPTIMIZATION FOR MASKING NETWORKS 366 TOWARD SCALABLE AND EFFICIENT SPEECH ENHANCEMENT

Sunwoo Kim, Minje Kim, Indiana University, United States of America

AUD-11.2: HGCN: HARMONIC GATED COMPENSATION NETWORK FOR SPEECH 371 ENHANCEMENT

Tianrui Wang, Weibin Zhu, Beijing Jiaotong University, China; Yingying Gao, Junlan Feng, Shilei Zhang, China Mobile Research Institute, China

AUD-11.3: SPEECH ENHANCEMENT WITH NEURAL HOMOMORPHIC SYNTHESIS 376

Wenbin Jiang, Zhijun Liu, Kai Yu, Fei Wen, Shanghai Jiao Tong University, China

AUD-11.4: A BAYESIAN PERMUTATION TRAINING DEEP REPRESENTATION 381 LEARNING METHOD FOR SPEECH ENHANCEMENT WITH VARIATIONAL AUTOENCODER

Yang Xiang, Aalborg University & Capturi A/S, Denmark; Jesper Lisby Højvang, Morten Højfeldt Rasmussen, Capturi A/S, Denmark; Mads Græsbøll Christensen, Aalborg University, Denmark

AUD-11.5: INTEGRATING STATISTICAL UNCERTAINTY INTO NEURAL 386 NETWORK-BASED SPEECH ENHANCEMENT

Huajian Fang, Tal Peer, Stefan Wermter, Timo Gerkmann, Universität Hamburg, Germany

AUD-11.6: UNSUPERVISED SPEECH ENHANCEMENT WITH SPEECH 391 RECOGNITION EMBEDDING AND DISENTANGLEMENT LOSSES

Viet Anh Trinh, The City University of New York, United States of America; Sebastian Braun, Microsoft, United States of America

AUD-12: MUSICAL EVENT DETECTION AND STRUCTURE ANALYSIS

AUD-12.1: MUSICYOLO: A SIGHT-SINGING ONSET/OFFSET DETECTION 396 FRAMEWORK BASED ON OBJECT DETECTION INSTEAD OF SPECTRUM FRAMES

Xianke Wang, Wei Xu, Weiming Yang, Wenqing Cheng, Huazhong University of Science and Technology, China

AUD-12.2: MODELING BEATS AND DOWNBEATS WITH A TIME-FREQUENCY TRANSFORMER	401
<i>Yun-Ning Hung, Georgia Institute of Technology, United States of America; Ju-Chiang Wang, Xuchen Song, Wei-Tsung Lu, Minz Won, ByteDance, United States of America</i>	
AUD-12.3: HIERARCHICAL CLASSIFICATION OF SINGING ACTIVITY, GENDER, AND TYPE IN COMPLEX MUSIC RECORDINGS	406
<i>Michael Krause, Meinard Müller, International Audio Laboratories Erlangen, Germany</i>	
AUD-12.4: DEEPCHORUS: A HYBRID MODEL OF MULTI-SCALE CONVOLUTION AND SELF-ATTENTION FOR CHORUS DETECTION	411
<i>Qiqi He, Xiaoheng Sun, Wei Li, Fudan University, School of Computer Science and Technology, China, China; Yi Yu, National Institute of Informatics (NII), Tokyo, Japan, Japan</i>	
AUD-12.5: TO CATCH A CHORUS, VERSE, INTRO, OR ANYTHING ELSE: ANALYZING A SONG WITH STRUCTURAL FUNCTIONS	416
<i>Ju-Chiang Wang, Jordan B. L. Smith, ByteDance, United States of America; Yun-Ning Hung, Georgia Institute of Technology, United States of America</i>	
AUD-12.6: A NOVEL 1D STATE SPACE FOR EFFICIENT MUSIC RHYTHMIC ANALYSIS	421
<i>Mojtaba Heydari, Zhiyao Duan, University of Rochester, United States of America; Matthew McCallum, Andreas Ehmman, Pandora Media, Inc., United States of America</i>	
AUD-13: SPATIAL AUDIO I	
AUD-13.1: UPMIXING VIA STYLE TRANSFER: A VARIATIONAL AUTOENCODER FOR DISENTANGLING SPATIAL IMAGES AND MUSICAL CONTENT	426
<i>Haici Yang, Minje Kim, Indiana University, United States of America; Sanna Wager, Spencer Russell, Mike Luo, Wontak Kim, Amazon Lab 126, United States of America</i>	
AUD-13.2: SPATIAL MIXUP: DIRECTIONAL LOUDNESS MODIFICATION AS DATA AUGMENTATION FOR SOUND EVENT LOCALIZATION AND DETECTION	431
<i>Ricardo Falcon Perez, Aalto University, Finland; Kazuki Shimada, Yuichiro Koyama, Shusuke Takahashi, Yuki Mitsufuji, Sony Group Corporation, Japan</i>	
AUD-13.3: TOWARDS FASTER CONTINUOUS MULTI-CHANNEL HRTF MEASUREMENTS BASED ON LEARNING SYSTEM MODELS	436
<i>Tobias Kabzinski, Peter Jax, RWTH Aachen University, Germany</i>	
AUD-13.4: TOWARDS FAST AND CONVENIENT END-TO-END HRTF PERSONALIZATION	441
<i>Bowen Zhi, Dmitry Zotkin, Ramani Duraiswami, VisiSonics Corporation, United States of America</i>	
AUD-13.6: WISHART LOCALIZATION PRIOR ON SPATIAL COVARIANCE MATRIX IN AMBISONIC SOURCE SEPARATION USING NON-NEGATIVE TENSOR FACTORIZATION	446
<i>Mateusz Guzik, Konrad Kowalczyk, AGH University of Science and Technology, Poland</i>	
AUD-14: MUSIC APPLICATIONS	
AUD-14.1: IMPROVING LYRICS ALIGNMENT THROUGH JOINT PITCH DETECTION	451
<i>Jiawen Huang, Emmanouil Benetos, Queen Mary University of London, United Kingdom of Great Britain and Northern Ireland; Sebastian Ewert, Spotify, Germany</i>	

AUD-14.2: LEARNING MUSIC AUDIO REPRESENTATIONS VIA WEAK LANGUAGE SUPERVISION	456
<i>Ilaria Manco, Emmanouil Benetos, Gyorgy Fazekas, Queen Mary University of London, United Kingdom of Great Britain and Northern Ireland; Elio Quinton, Universal Music Group, United Kingdom of Great Britain and Northern Ireland</i>	
AUD-14.3: ON THE PREDICTION OF THE FREQUENCY RESPONSE OF A WOODEN PLATE FROM ITS MECHANICAL PARAMETERS	461
<i>David Giuseppe Badiane, Raffaele Malvermi, Sebastian Gonzalez, Fabio Antonacci, Augusto Sarti, Politecnico di Milano, Italy</i>	
AUD-14.5: AUTOMATIC DJ TRANSITIONS WITH DIFFERENTIABLE AUDIO EFFECTS AND GENERATIVE ADVERSARIAL NETWORKS	466
<i>Bo-Yu Chen, Wei-Han Hsu, Yi-Hsuan Yang, Academia Sinica, Taiwan; Wei-Hsiang Liao, Marco Antonio Martinez Ramirez, Yuki Mitsufuji, Sony Group Corporation, Japan</i>	
AUD-15: DETECTION AND CLASSIFICATION OF ACOUSTIC SCENES AND EVENTS IV: ON-DEVICE CONSIDERATIONS	
AUD-15.1: SELF-SUPERVISED REPRESENTATION LEARNING FOR UNSUPERVISED ANOMALOUS SOUND DETECTION UNDER DOMAIN SHIFT	471
<i>Han Chen, Yan Song, Li-Rong Dai, Ian McLoughlin, University of Science and Technology of China, China; Lin Liu, iFLYTEK Research, iFLYTEK CO., LTD, China</i>	
AUD-15.3: FEDERATED SELF-TRAINING FOR DATA-EFFICIENT AUDIO RECOGNITION	476
<i>Vasileios Tsouvalas, Tanir Ozcelebi, Eindhoven University of Technology, Netherlands; Aaqib Saeed, Philips Research, Netherlands</i>	
AUD-15.4: FEDERATED SELF-SUPERVISED LEARNING FOR ACOUSTIC EVENT CLASSIFICATION	481
<i>Meng Feng, Massachusetts Institute of Technology, United States of America; Chieh-Chi Kao, Qingming Tang, Ming Sun, Viktor Rozgic, Spyros Matsoukas, Chao Wang, Amazon, United States of America</i>	
AUD-15.5: TEMPORAL KNOWLEDGE DISTILLATION FOR ON-DEVICE AUDIO CLASSIFICATION	486
<i>Kwanghee Choi, Martin Kersner, Jacob Morton, Buru Chang, Hyperconnect, Korea, Republic of</i>	
AUD-15.6: STREAMING ON-DEVICE DETECTION OF DEVICE DIRECTED SPEECH FROM VOICE AND TOUCH-BASED INVOCATION	491
<i>Ognjen (Oggi) Rudovic, Akanksha Bindal, Vineet Garg, Pramod Simha, Pranay Dighe, Sachin Kajarekar, Apple Inc., United States of America</i>	
AUD-16: MULTICHANNEL BLIND SOURCE SEPARATION: LATENT VARIABLE ANALYSIS	
AUD-16.1: MULTI-FRAME FULL-RANK SPATIAL COVARIANCE ANALYSIS FOR UNDERDETERMINED BSS IN REVERBERANT ENVIRONMENTS	496
<i>Hiroshi Sawada, Rintaro Ikeshita, Keisuke Kinoshita, Tomohiro Nakatani, NTT Corporation, Japan</i>	
AUD-16.2: FLOW-BASED FAST MULTICHANNEL NONNEGATIVE MATRIX FACTORIZATION FOR BLIND SOURCE SEPARATION	501
<i>Aditya Arie Nugraha, Kouhei Sekiguchi, RIKEN, Japan; Mathieu Fontaine, Télécom Paris, France; Yoshiaki Bando, National Institute of Advanced Industrial Science and Technology (AIST), Japan; Kazuyoshi Yoshii, Kyoto University, Japan</i>	

AUD-16.3: HARVESTING PARTIALLY-DISJOINT TIME-FREQUENCY INFORMATION	506
FOR IMPROVING DEGENERATE UNMIXING ESTIMATION TECHNIQUE	
<i>Yudong He, Qifeng Chen, Richard H.Y. So, The Hong Kong University of Science and Technology, Hong Kong; He Wang, HKUST-Shenzhen Research Institute, China</i>	
AUD-16.4: INVESTIGATION AND COMPARISON OF OPTIMIZATION METHODS FOR	511
VARIATIONAL AUTOENCODER-BASED UNDERDETERMINED MULTICHANNEL SOURCE SEPARATION	
<i>Shogo Seki, Hirokazu Kameoka, NTT Corporation, Japan; Li Li, Nagoya University, Japan</i>	
AUD-16.6: HBP: AN EFFICIENT BLOCK PERMUTATION SOLVER USING	516
HUNGARIAN ALGORITHM AND SPECTROGRAM INPAINTING FOR MULTICHANNEL AUDIO SOURCE SEPARATION	
<i>Li Li, Nagoya University, Japan; Hirokazu Kameoka, Shogo Seki, NTT Communication Science Laboratories, NTT Corporation, Japan</i>	
 AUD-17: SPEECH SEPARATION I	
AUD-17.1: EAD-CONFORMER: A CONFORMER-BASED	521
ENCODER-ATTENTION-DECODER-NETWORK FOR MULTI-TASK AUDIO SOURCE SEPARATION	
<i>Chenxing Li, Yang Wang, Feng Deng, Zhuo Zhang, Xiaorui Wang, Zhongyuan Wang, Kuai Shou, China</i>	
AUD-17.2: THE COCKTAIL FORK PROBLEM: THREE-STEM AUDIO SEPARATION	526
FOR REAL-WORLD SOUNDTRACKS	
<i>Darius Petermann, Indiana University, United States of America; Gordon Wichern, Jonathan Le Roux, Mitsubishi Electric Research Laboratories (MERL), United States of America; Zhong-Qiu Wang, Carnegie Mellon University, United States of America</i>	
AUD-17.3: PHASE SHIFTED BEDROSIAN FILTERBANK: AN INTERPRETABLE AUDIO	531
FRONT-END FOR TIME-DOMAIN AUDIO SOURCE SEPARATION	
<i>Félix Mathieu, Gael Richard, Geoffroy Peeters, Telecom Paris, France; Thomas Courtat, Thales, France</i>	
AUD-17.4: HARMONICITY PLAYS A CRITICAL ROLE IN DNN BASED VERSUS IN	536
BIOLOGICALLY-INSPIRED MONAURAL SPEECH SEGREGATION SYSTEMS	
<i>Rahil Parikh, Carol Espy-Wilson, Shihab Shamma, University of Maryland College Park, United States of America; Ilya Kavalero, Google Inc., United States of America</i>	
AUD-17.5: MULTI-CHANNEL NARROW-BAND DEEP SPEECH SEPARATION WITH	541
FULL-BAND PERMUTATION INVARIANT TRAINING	
<i>Changsheng Quan, Zhejiang University, China; Xiaofei Li, Westlake University, China</i>	
 AUD-18: MEDICAL ACOUSTICS	
AUD-18.1: CSENET: COMPLEX SQUEEZE-AND-EXCITATION NETWORK FOR	546
SPEECH DEPRESSION LEVEL PREDICTION	
<i>Cunhang Fan, Zhao Lv, Shengbing Pei, Anhui Province Key Laboratory of Multimodal Cognitive Computation, School of Computer Science and Technology, Anhui University, Hefei 230601, China, China; Mingyue Niu, School of Computer and Information Engineering, Tianjin Normal University, Tianjin 300387, China, China</i>	
AUD-18.2: UBILUNG: MULTI-MODAL PASSIVE-BASED LUNG HEALTH ASSESSMENT.....	551
<i>Ebrahim Nemati, Viswam Nathan, Korosh Vatanparvar, Tousif Ahmed, Md. Mahbubur Rahman, Dan McCaffrey, Jilong Kuang, Alex Gao, Samsung Research America, United States of America; Xuhai Xu, University of Washington, United States of America</i>	
AUD-18.3: THE SECOND DICOVA CHALLENGE: DATASET AND PERFORMANCE	556
ANALYSIS FOR DIAGNOSIS OF COVID-19 USING ACOUSTICS	
<i>Neeraj Kumar Sharma, Srikanth Raj Chetupalli, Debarpan Bhattacharya, Debottam Dutta, Pravin Mote, Sriram Ganapathy, Indian Institute of Science, India</i>	

AUD-18.4: SUPERVISED AND SELF-SUPERVISED PRETRAINING BASED COVID-19 DETECTION USING ACOUSTIC BREATHING/COUGH/SPEECH SIGNALS	561
<i>Xing-Yu Chen, Qiu-Shi Zhu, Jie Zhang, Li-Rong Dai, University of Science and Technology of China, China</i>	
AUD-18.5: EXPLORING AUDITORY ACOUSTIC FEATURES FOR THE DIAGNOSIS OF COVID-19	566
<i>Madhu Kamble, Jose Patino, Maria A. Zuluaga, Massimiliano Todisco, EURECOM, France, France</i>	
AUD-19: MODELING, ANALYSIS, AND SYNTHESIS OF ACOUSTIC ENVIRONMENTS	
AUD-19.1: FAST-RIR: FAST NEURAL DIFFUSE ROOM IMPULSE RESPONSE GENERATOR	571
<i>Anton Jeran Ratnarajah, Zhenyu Tang, Dinesh Manocha, University of Maryland, College Park, United States of America; Shi-Xiong Zhang, Meng Yu, Dong Yu, Tencent AI Lab, United States of America</i>	
AUD-19.2: REGION-TO-REGION KERNEL INTERPOLATION OF ACOUSTIC TRANSFER FUNCTION WITH DIRECTIONAL WEIGHTING	576
<i>Juliano Garcia do Carmo Ribeiro, Shoichi Koyama, Hiroshi Saruwatari, The University of Tokyo, Japan</i>	
AUD-19.3: BLIND REVERBERATION TIME ESTIMATION IN DYNAMIC ACOUSTIC CONDITIONS	581
<i>Philipp Götz, Emanuël A. P. Habets, International Audio Laboratories Erlangen, Germany; Cagdas Tuna, Andreas Walther, Fraunhofer Institute for Integrated Circuits IIS, Germany</i>	
AUD-19.4: SPARSE MODELING OF THE EARLY PART OF NOISY ROOM IMPULSE RESPONSES WITH SPARSE BAYESIAN LEARNING	586
<i>Maozhong Fu, Yuhang Li, Xiamen University, China; Jesper Rindom Jensen, Mads Græsbøll Christensen, Aalborg University, Denmark</i>	
AUD-19.5: IMPROVED SIMULATION OF REALISTICALLY-SPATIALISED SIMULTANEOUS SPEECH USING MULTI-CAMERA ANALYSIS IN THE CHIME-5 DATASET	591
<i>Jack Deadman, Jon Barker, University of Sheffield, United Kingdom of Great Britain and Northern Ireland</i>	
AUD-19.6: A DATA-DRIVEN APPROACH FOR ACOUSTIC PARAMETER SIMILARITY ESTIMATION OF SPEECH RECORDING	596
<i>Mattia Papa, Clara Borrelli, Paolo Bestagini, Fabio Antonacci, Augusto Sarti, Stefano Tubaro, Politecnico di Milano, Italy</i>	
AUD-20: MUSIC CLASSIFICATION, IDENTIFICATION, AND MELODY EXTRACTION	
AUD-20.1: VIOLINIST IDENTIFICATION USING NOTE-LEVEL TIMBRE FEATURE DISTRIBUTIONS	601
<i>Yudong Zhao, György Fazekas, Mark Sandler, Queen Mary University of London, United Kingdom of Great Britain and Northern Ireland</i>	
AUD-20.2: S3T: SELF-SUPERVISED PRE-TRAINING WITH SWIN TRANSFORMER FOR MUSIC CLASSIFICATION	606
<i>Hang Zhao, Bilei Zhu, Zejun Ma, ByteDance, China; Chen Zhang, Kejun Zhang, Zhejiang University, China</i>	
AUD-20.3: AMBIGUITY MODELLING WITH LABEL DISTRIBUTION LEARNING FOR MUSIC CLASSIFICATION	611
<i>Morgan Buisson, Pablo Alonso Jimenez, Dmitry Bogdanov, Universitat Pompeu Fabra, Spain</i>	
AUD-20.4: BYTECOVER2: TOWARDS DIMENSIONALITY REDUCTION OF LATENT EMBEDDING FOR EFFICIENT COVER SONG IDENTIFICATION	616
<i>Xingjian Du, Zijie Wang, Bilei Zhu, Zejun Ma, Bytedance AI Lab, China; Ke Chen, University of California San Diego, China</i>	

AUD-20.5: TONET: TONE-OCTAVE NETWORK FOR SINGING MELODY EXTRACTION FROM POLYPHONIC MUSIC	621
<i>Ke Chen, Taylor Berg-Kirkpatrick, Shlomo Dubnov, University of California San Diego, United States of America; Shuai Yu, Wei Li, Fudan University, China; Cheng-i Wang, Smule Inc., United States of America</i>	
AUD-20.6: HIERARCHICAL GRAPH-BASED NEURAL NETWORK FOR SINGING MELODY EXTRACTION	626
<i>Shuai Yu, Xi Chen, Wei Li, Fudan University, China</i>	
AUD-21: DETECTION AND CLASSIFICATION OF ACOUSTIC SCENES AND EVENTS V: CLASSIFICATION	
AUD-21.1: ON THE IMPACT OF NORMALIZATION STRATEGIES IN UNSUPERVISED ADVERSARIAL DOMAIN ADAPTATION FOR ACOUSTIC SCENE CLASSIFICATION	631
<i>Michel Olvera, Emmanuel Vincent, Université de Lorraine, CNRS, Inria, Loria, France; Gilles Gasso, LITIS, Université & INSA Rouen Normandie, France</i>	
AUD-21.2: IMPROVING BIRD CLASSIFICATION WITH UNSUPERVISED SOUND SEPARATION	636
<i>Tom Denton, Google, United States of America; Scott Wisdom, John R. Hershey, Google Research, United States of America</i>	
AUD-21.3: SCALABLE NEURAL ARCHITECTURES FOR END-TO-END ENVIRONMENTAL SOUND CLASSIFICATION	641
<i>Francesco Paissan, Alberto Ancilotto, Alessio Brutti, Elisabetta Farella, Fondazione Bruno Kessler, Italy</i>	
AUD-21.4: HTS-AT: A HIERARCHICAL TOKEN-SEMANTIC AUDIO TRANSFORMER FOR SOUND CLASSIFICATION AND DETECTION	646
<i>Ke Chen, Taylor Berg-Kirkpatrick, Shlomo Dubnov, University of California San Diego, United States of America; Xingjian Du, Bilei Zhu, Zejun Ma, Bytedance Inc., China</i>	
AUD-21.5: HYBRID ATTENTION-BASED PROTOTYPICAL NETWORKS FOR FEW-SHOT SOUND CLASSIFICATION	651
<i>You Wang, David Anderson, Georgia Institute of Technology, United States of America</i>	
AUD-22: ACTIVE NOISE CONTROL, ECHO REDUCTION, AND FEEDBACK REDUCTION I: NOISE ROBUSTNESS	
AUD-22.1: END-TO-END COMPLEX-VALUED MULTIDILATED CONVOLUTIONAL NEURAL NETWORK FOR JOINT ACOUSTIC ECHO CANCELLATION AND NOISE SUPPRESSION	656
<i>Karn N. Watcharasupat, Thi Ngoc Tho Nguyen, Woon-Seng Gan, Nanyang Technological University, Singapore; Shengkui Zhao, Bin Ma, Alibaba Group, Singapore</i>	
AUD-22.2: NN3A: NEURAL NETWORK SUPPORTED ACOUSTIC ECHO CANCELLATION, NOISE SUPPRESSION AND AUTOMATIC GAIN CONTROL FOR REAL-TIME COMMUNICATIONS	661
<i>Ziteng Wang, Yueyue Na, Biao Tian, Qiang Fu, Alibaba Group, China</i>	
AUD-22.3: DEEP RESIDUAL ECHO SUPPRESSION AND NOISE REDUCTION: A MULTI-INPUT FCRN APPROACH IN A HYBRID SPEECH ENHANCEMENT SYSTEM	666
<i>Jan Franzen, Tim Fingscheidt, Technische Universität Braunschweig, Germany</i>	
AUD-22.4: NEURAL CASCADE ARCHITECTURE FOR JOINT ACOUSTIC ECHO AND NOISE SUPPRESSION	671
<i>Hao Zhang, DeLiang Wang, The Ohio State University, United States of America</i>	

AUD-22.5: CASCADE MULTI-CHANNEL NOISE REDUCTION AND ACOUSTIC FEEDBACK CANCELLATION	676
<i>Santiago Ruiz, Toon van Waterschoot, Marc Moonen, KU Leuven, Belgium</i>	
 AUD-23: SPEECH SEPARATION II	
AUD-23.1: SKIM: SKIPPING MEMORY LSTM FOR LOW-LATENCY REAL-TIME CONTINUOUS SPEECH SEPARATION	681
<i>Chenda Li, Yanmin Qian, Shanghai Jiao Tong University, China; Lei Yang, Weiqin Wang, Samsung Research China - Beijing, China</i>	
AUD-23.2: ADAPTING SPEECH SEPARATION TO REAL-WORLD MEETINGS USING MIXTURE INVARIANT TRAINING	686
<i>Aswin Sivaraman, Indiana University, United States of America; Scott Wisdom, Hakan Erdogan, John R. Hershey, Google Research, United States of America</i>	
AUD-23.3: QUANTIFYING DISCRIMINABILITY BETWEEN NMF BASES.....	691
<i>Eisuke Konno, Daisuke Saito, Nobuaki Minematsu, The University of Tokyo, Japan</i>	
AUD-23.4: LOCATION-BASED TRAINING FOR MULTI-CHANNEL TALKER-INDEPENDENT SPEAKER SEPARATION	696
<i>Hassan Taherian, Ke Tan, DeLiang Wang, The Ohio State University, United States of America</i>	
AUD-23.5: SDR — MEDIUM RARE WITH FAST COMPUTATIONS.....	701
<i>Robin Scheibler, LINE Corporation, Japan</i>	
AUD-23.6: ATTENTIONPIT: SOFT PERMUTATION INVARIANT TRAINING FOR AUDIO SOURCE SEPARATION WITH ATTENTION MECHANISM	706
<i>Hirokazu Kameoka, Shogo Seki, Li Li, Chihiro Watanabe, Nippon Telegraph and Telephone Corporation, Japan</i>	
 AUD-24: ACOUSTIC SENSOR ARRAY PROCESSING I: LOCALIZATION	
AUD-24.1: LOCATE THIS, NOT THAT: CLASS-CONDITIONED SOUND EVENT DOA ESTIMATION	711
<i>Olga Slizovskaia, Gordon Wichern, Zhong-Qiu Wang, Jonathan Le Roux, Mitsubishi Electric Research Laboratories (MERL), United States of America</i>	
AUD-24.2: SALSA-LITE: A FAST AND EFFECTIVE FEATURE FOR POLYPHONIC SOUND EVENT LOCALIZATION AND DETECTION WITH MICROPHONE ARRAYS	716
<i>Thi Ngoc Tho Nguyen, Karn N. Watcharasupat, Woon-Seng Gan, Nanyang Technological University, Singapore; Douglas L. Jones, University of Illinois Urbana-Champaign, United States of America; Huy Phan, Queen Mary University of London, United Kingdom of Great Britain and Northern Ireland</i>	
AUD-24.3: SRP-DNN: LEARNING DIRECT-PATH PHASE DIFFERENCE FOR MULTIPLE MOVING SOUND SOURCE LOCALIZATION	721
<i>Bing Yang, Hong Liu, Peking University, China; Xiaofei Li, Westlake University & Westlake Institute for Advanced Study, China</i>	
AUD-24.4: CLOSED-FORM SINGLE SOURCE DIRECTION-OF-ARRIVAL ESTIMATOR USING FIRST-ORDER RELATIVE HARMONIC COEFFICIENTS	726
<i>Yonggang Hu, Australian National University, Australia; Sharon Gannot, Bar-Ilan University, Israel</i>	
AUD-24.5: A SLIDE-SAVE BASED FRAMEWORK FOR MULTI-SOURCE DOA EXTRACTION WITH SPATIAL CLOSELY SEPARATED SOURCES	731
<i>Jianhua Geng, Sifan Wang, Xin Lou, ShanghaiTech University, China</i>	

AUD-24.6: AN END-TO-END DEEP LEARNING FRAMEWORK FOR MULTIPLE AUDIO SOURCE SEPARATION AND LOCALIZATION	736
<i>Yu Chen, Bowen Liu, Zijian Zhang, Hun-Seok Kim, University of Michigan, United States of America</i>	
 AUD-25: ACTIVE NOISE CONTROL, ECHO REDUCTION, AND FEEDBACK REDUCTION II	
AUD-25.1: DEEP ADAPTATION CONTROL FOR ACOUSTIC ECHO CANCELLATION	741
<i>Amir Ivry, Israel Cohen, Baruch Berdugo, Technion - Israel Institute of Technology, Israel</i>	
AUD-25.2: OFF-THE-SHELF DEEP INTEGRATION FOR RESIDUAL-ECHO SUPPRESSION	746
<i>Amir Ivry, Israel Cohen, Baruch Berdugo, Technion - Israel Institute of Technology, Israel</i>	
AUD-25.3: A COMPLEX SPECTRAL MAPPING WITH INPLACE CONVOLUTION RECURRENT NEURAL NETWORKS FOR ACOUSTIC ECHO CANCELLATION	751
<i>Chenggang Zhang, Jinjiang Liu, Xueliang Zhang, Inner Mongolia University, China</i>	
AUD-25.4: DEEP ADAPTIVE AEC: HYBRID OF DEEP LEARNING AND ADAPTIVE ACOUSTIC ECHO CANCELLATION	756
<i>Hao Zhang, The Ohio State University, United States of America; Srivatsan Kandadai, Harsha Rao, Tarun Pruthi, Trausti Kristjansson, Amazon Inc, United States of America; Minje Kim, Indiana University, United States of America</i>	
AUD-25.5: COMPUTATIONALLY EFFICIENT FIXED-FILTER ANC FOR SPEECH BASED ON LONG-TERM PREDICTION FOR HEADPHONE APPLICATIONS	761
<i>Yurii Iotov, Mads Græsbøll Christensen, Aalborg University, Denmark; Sidsel Marie Nørholm, Valiantsin Belyi, Mads Dyrholm, GN Audio A/S, Denmark</i>	
AUD-25.6: END-TO-END DEEP LEARNING-BASED ADAPTATION CONTROL FOR FREQUENCY-DOMAIN ADAPTIVE SYSTEM IDENTIFICATION	766
<i>Thomas Haubner, Andreas Brendel, Walter Kellermann, Friedrich-Alexander-University Erlangen-Nuremberg, Germany</i>	
 AUD-26: MUSIC AND LYRICS TRANSCRIPTION	
AUD-26.1: A FEW-SAMPLE STRATEGY FOR GUITAR TABLATURE TRANSCRIPTION BASED ON INHARMONICITY ANALYSIS AND PLAYABILITY CONSTRAINTS	771
<i>Grigoris Bastas, Maximos Kaliakatsos-Papakostas, Vassilis Katsouros, Athena Research Center, Greece; Stefanos Koutoupis, Petros Maragos, National Technical University of Athens, Greece</i>	
AUD-26.2: EXPLORING TRANSFORMER’S POTENTIAL ON AUTOMATIC PIANO TRANSCRIPTION	776
<i>Longshen Ou, Ziyi Guo, Ye Wang, National University of Singapore, Singapore; Emmanouil Benetos, Queen Mary University of London, United Kingdom of Great Britain and Northern Ireland; Jiqing Han, Harbin Institute of Technology, China</i>	
AUD-26.3: A LIGHTWEIGHT INSTRUMENT-AGNOSTIC MODEL FOR POLYPHONIC NOTE TRANSCRIPTION AND MULTIPITCH ESTIMATION	781
<i>Rachel Bittner, Juan José Bosch, David Rubinstein, Sebastian Ewert, Spotify, France; Gabriel Meseguer Brocal, Institut de recherche et coordination acoustique/musique (IRCAM), France</i>	
AUD-26.4: TOWARDS AUTOMATIC TRANSCRIPTION OF POLYPHONIC ELECTRIC GUITAR MUSIC: A NEW DATASET AND A MULTI-LOSS TRANSFORMER MODEL	786
<i>Yu-Hua Chen, Jyh-Shing Roger Jang, National Taiwan University, Taiwan; Wen-Yi Hsiao, Tsu-Kuang Hsieh, Yi-Hsuan Yang, Taiwan AI Labs, Taiwan</i>	
AUD-26.5: GENRE-CONDITIONED ACOUSTIC MODELS FOR AUTOMATIC LYRICS TRANSCRIPTION OF POLYPHONIC MUSIC	791
<i>Xiaoxue Gao, Chitraksha Gupta, Haizhou Li, National University of Singapore, Singapore</i>	

AUD-26.6: PSEUDO-LABEL TRANSFER FROM FRAME-LEVEL TO NOTE-LEVEL IN A TEACHER-STUDENT FRAMEWORK FOR SINGING TRANSCRIPTION FROM POLYPHONIC MUSIC	796
<i>Sangeun Kum, Jongpil Lee, Keunhyoung Luke Kim, Taehyoung Kim, Neutune Research, Korea, Republic of; Juhan Nam, KAIST, Korea, Republic of</i>	

AUD-27: DETECTION AND CLASSIFICATION OF ACOUSTIC SCENES AND EVENTS VI: EVENTS

AUD-27.1: SOUND EVENT DETECTION GUIDED BY SEMANTIC CONTEXTS OF SCENES	801
<i>Noriyuki Tonami, Ryotaro Nagase, Yuki Okamoto, Takahiro Fukumori, Yoichi Yamashita, Ritsumeikan University, Japan; Keisuke Imoto, Doshisha University, Japan</i>	

AUD-27.2: CNN-TRANSFORMER WITH SELF-ATTENTION NETWORK FOR SOUND EVENT DETECTION	806
<i>Keigo Wakayama, Shoichiro Saito, NTT Corporation, Japan</i>	

AUD-27.3: A MUTUAL LEARNING FRAMEWORK FOR FEW-SHOT SOUND EVENT DETECTION	811
<i>Dongchao Yang, Helin Wang, Yuexian Zou, Zhongjie Ye, Peking University, China; Wenwu Wang, University of Surrey, United Kingdom of Great Britain and Northern Ireland</i>	

AUD-27.4: ANOMALOUS SOUND DETECTION USING SPECTRAL-TEMPORAL INFORMATION FUSION	816
<i>Youde Liu, Jian Guan, Harbin Engineering University, China; Qiaoxi Zhu, University of Technology Sydney, Australia; Wenwu Wang, University of Surrey, United Kingdom of Great Britain and Northern Ireland</i>	

AUD-27.5: SPARSE SELF-ATTENTION FOR SEMI-SUPERVISED SOUND EVENT DETECTION	821
<i>Yadong Guan, Jiabin Xue, Guibin Zheng, Jiqing Han, Harbin Institute of Technology, China</i>	

AUD-27.6: PEER COLLABORATIVE LEARNING FOR POLYPHONIC SOUND EVENT DETECTION	826
<i>Hayato Endo, Hiromitsu Nishizaki, University of Yamanashi, Japan</i>	

AUD-28: SIGNAL ENHANCEMENT AND RESTORATION II

AUD-28.1: POSTGAN: A GAN-BASED POST-PROCESSOR TO ENHANCE THE QUALITY OF CODED SPEECH	831
<i>Srikanth Korse, Nicola Pia, Kishan Gupta, Guillaume Fuchs, Fraunhofer IIS, Erlangen, Germany, Germany</i>	

AUD-28.2: A DNN BASED POST-FILTER TO ENHANCE THE QUALITY OF CODED SPEECH IN MDCT DOMAIN	836
<i>Kishan Gupta, Srikanth Korse, Bernd Edler, Guillaume Fuchs, Fraunhofer IIS, Germany</i>	

AUD-28.3: A TWO-STAGE U-NET FOR HIGH-FIDELITY DENOISING OF HISTORICAL RECORDINGS	841
<i>Eloi Moliner, Vesa Välimäki, Aalto University, Finland</i>	

AUD-28.4: EXPERTS VERSUS ALL-ROUNDERS: TARGET LANGUAGE EXTRACTION FOR MULTIPLE TARGET LANGUAGES	846
<i>Marvin Borsdorf, Kevin Scheck, Tanja Schultz, University of Bremen, Germany; Haizhou Li, National University of Singapore, Singapore</i>	

AUD-28.5: CATEGORY-ADAPTED SOUND EVENT ENHANCEMENT WITH WEAKLY LABELED DATA	851
<i>Guangwei Li, Xuenan Xu, Mengyue Wu, Kai Yu, Shanghai Jiao Tong University, China; Heinrich Dinkel, Xiaomi Corporation, China</i>	

AUD-28.6: SEQUENTIAL MCMC METHODS FOR AUDIO SIGNAL ENHANCEMENT	856
<i>Ruben Claveria, Simon Godsill, University of Cambridge, United Kingdom of Great Britain and Northern Ireland</i>	
 AUD-29: AUDIO AND SPEECH CODING	
AUD-29.3: ARCHITECTURE FOR VARIABLE BITRATE NEURAL SPEECH CODEC	861
WITH CONFIGURABLE COMPUTATION COMPLEXITY	
<i>Tejas Jayashankar, Massachusetts Institute of Technology, United States of America; Thilo Koehler, Kaustubh Kalgaonkar, Zhiping Xiu, Jilong Wu, Ju Lin, Prabhav Agrawal, Qing He, Facebook AI, United States of America</i>	
AUD-29.4: END-TO-END NEURAL SPEECH CODING FOR REAL-TIME	866
COMMUNICATIONS	
<i>Xue Jiang, Chengyu Zheng, Yuan Zhang, Communication University of China, China; Xiulian Peng, Huaying Xue, Yan Lu, Microsoft Research Asia, China</i>	
AUD-29.5: DEEP NEURAL NETWORK (DNN) AUDIO CODER USING A	871
PERCEPTUALLY IMPROVED TRAINING METHOD	
<i>Seungmin Shin, Joon Byun, Youngcheol Park, Intelligent Signal Processing Lab, Yonsei University, Wonju, Korea, Republic of; Jongmo Sung, Seungkwon Beack, Electronics and Telecommunications Research Institute (ETRI), Daejeon, Korea, Republic of</i>	
AUD-29.6: PROGRESSIVE MULTI-STAGE NEURAL AUDIO CODING WITH GUIDED	876
REFERENCES	
<i>Chanwoo Lee, Hyungseob Lim, Jihyun Lee, Hong-Goo Kang, Yonsei university, Korea, Republic of; Inseon Jang, Electronics and Telecommunications Research Institution, Korea, Republic of</i>	
 AUD-30: AUDIO QUALITY AND SPEECH INTELLIGIBILITY MEASURES	
AUD-30.1: VOCBENCH: A NEURAL VOCODER BENCHMARK FOR SPEECH	881
SYNTHESIS	
<i>Ehab A. AlBadawy, Ming-Ching Chang, University at Albany, State University of New York, United States of America; Andrew Gibiansky, Qing He, Jilong Wu, Facebook AI, United States of America; Siwei Lyu, University at Buffalo, State University of New York, USA, United States of America</i>	
AUD-30.2: DNSMOS P.835: A NON-INTRUSIVE PERCEPTUAL OBJECTIVE SPEECH	886
QUALITY METRIC TO EVALUATE NOISE SUPPRESSORS	
<i>Chandan Reddy, Google, United States of America; Vishak Gopal, Ross Cutler, Microsoft, United States of America</i>	
AUD-30.3: SQAPP: NO-REFERENCE SPEECH QUALITY ASSESSMENT VIA PAIRWISE	891
PREFERENCE	
<i>Pranay Manocha, Adam Finkelstein, Princeton University, United States of America; Zeyu Jin, Adobe Research, United States of America</i>	
AUD-30.4: LDNET: UNIFIED LISTENER DEPENDENT MODELING IN MOS	896
PREDICTION FOR SYNTHETIC SPEECH	
<i>Wen-Chin Huang, Tomoki Toda, Nagoya University, Japan; Erica Cooper, Junichi Yamagishi, National Institute of Informatics, Japan</i>	
AUD-30.5: AECMOS: A SPEECH QUALITY ASSESSMENT METRIC FOR ECHO	901
IMPAIRMENT	
<i>Marju Purin, Sten Sootla, Mateja Sponza, Ando Saabas, Ross Cutler, Microsoft Corporation, Estonia</i>	
AUD-30.6: MOS PREDICTOR FOR SYNTHETIC SPEECH WITH I-VECTOR INPUTS.....	906
<i>Miao Liu, Jing Wang, Beijing Institute of Technology, China; Shicong Li, Fei Xiang, Xiaomi Inc., China; Yue Yao, Lidong Yang, Inner Mongolia University of Science and Technology, China</i>	

AUD-31: MULTICHANNEL PROCESSING

AUD-31.1: WAVE-DOMAIN APPROACH FOR CANCELLING NOISE ENTERING OPEN WINDOWS 911

Daan Ratering, Technical University Delft, Netherlands; W. Bastiaan Kleijn, Victoria University of Wellington, New Zealand; Jean Gonzalez Silva, Riccardo M. G. Ferrari, Delft University of Technology, Netherlands

AUD-31.2: ON SYNCHRONIZATION OF WIRELESS ACOUSTIC SENSOR NETWORKS 916 IN THE PRESENCE OF TIME-VARYING SAMPLING RATE OFFSETS AND SPEAKER CHANGES

Tobias Gburrek, Joerg Schmalenstroer, Reinhold Haeb-Umbach, Paderborn University, Germany

AUD-31.3: PICKNET: REAL-TIME CHANNEL SELECTION FOR AD HOC 921 MICROPHONE ARRAYS

Takuya Yoshioka, Xiaofei Wang, Dongmei Wang, Microsoft, United States of America

AUD-31.4: END-TO-END ALEXA DEVICE ARBITRATION 926

Jarred Barber, Tao Zhang, Amazon, United States of America; Yifeng Fan, University of Illinois, Urbana-Champaign, United States of America

AUD-31.5: INSTANTANEOUS LINEAR DIMENSIONALITY REDUCTION OF 931 MULTICHANNEL TIME-SERIES SIGNAL FOR ARRAY SIGNAL PROCESSING

Natsuki Ueno, Nobutaka Ono, Tokyo Metropolitan University, Japan

AUD-31.6: GENERALIZED TIME DOMAIN VELOCITY VECTOR 936

Srdan Kitić, Jérôme Daniel, Orange, France

AUD-32: AUDIO MUSIC SYNTHESIS

AUD-32.1: DIFFERENTIABLE DIGITAL SIGNAL PROCESSING MIXTURE MODEL 941 FOR SYNTHESIS PARAMETER EXTRACTION FROM MIXTURE OF HARMONIC SOUNDS

Masaya Kawamura, Tomohiko Nakamura, Hiroshi Saruwatari, The University of Tokyo, Japan; Daichi Kitamura, National Institute of Technology, Kagawa College, Japan; Yu Takahashi, Kazunobu Kondo, Yamaha Corporation, Japan

AUD-32.2: THE MIRRORNET : LEARNING AUDIO SYNTHESIZER CONTROLS 946 INSPIRED BY SENSORIMOTOR INTERACTION

Yashish M. Siriwardena, Shihab Shamma, Institute for Systems Research, University of Maryland College Park, United States of America; Guilhem Marion, Laboratoire des Systèmes Perceptifs, École Normale Supérieure, PSL University, France

AUD-32.3: DEEP PERFORMER: SCORE-TO-AUDIO MUSIC PERFORMANCE 951 SYNTHESIS

Hao-Wen Dong, Taylor Berg-Kirkpatrick, Julian McAuley, University of California San Diego, United States of America; Cong Zhou, Dolby Laboratories, United States of America

AUD-32.4: KARASINGER: SCORE-FREE SINGING VOICE SYNTHESIS WITH VQ-VAE 956 USING MEL-SPECTROGRAMS

Chien-Feng Liao, Jen-Yu Liu, Yi-Hsuan Yang, Taiwan AI Labs, Taiwan

AUD-32.5: ADVERSARIAL AUDIO SYNTHESIS USING A HARMONIC-PERCUSSIVE 961 DISCRIMINATOR

Jihyun Lee, Hyungseob Lim, Chanwoo Lee, Hong-Goo Kang, Yonsei University, Korea, Republic of; Inseon Jang, Electronics and Telecommunications Research Institution, Korea, Republic of

AUD-32.6: SLEEPGAN: TOWARDS PERSONALIZED SLEEP THERAPY MUSIC 966

Jing Yang, Swiss Federal Institute of Technology in Zurich, Switzerland; Chulhong Min, Akhil Mathur, Fahim Kawsar, Nokia Bell Labs, United Kingdom of Great Britain and Northern Ireland

AUD-33: EXTENDED EVALUATION AND CAPTIONING

AUD-33.1: DIVERSITY-CONTROLLABLE AND ACCURATE AUDIO CAPTIONING BASED ON NEURAL CONDITION 971

Xuenan Xu, Mengyue Wu, Kai Yu, Shanghai Jiao Tong University, China

AUD-33.2: AUDIOCLIP: EXTENDING CLIP TO IMAGE, TEXT AND AUDIO..... 976

Andrey Guzhov, Federico Raue, Jörn Hees, Andreas Dengel, Deutsches Forschungszentrum für Künstliche Intelligenz GmbH, Germany

AUD-33.3: CAN AUDIO CAPTIONS BE EVALUATED WITH IMAGE CAPTION METRICS? 981

Zelin Zhou, Zhiling Zhang, Xuenan Xu, Zeyu Xie, Mengyue Wu, Kenny Zhu, Shanghai Jiao Tong University, China

AUD-33.4: A DATA-DRIVEN COGNITIVE SALIENCE MODEL FOR OBJECTIVE PERCEPTUAL AUDIO QUALITY ASSESSMENT 986

Pablo M. Delgado, Jürgen Herre, International Audio Laboratories Erlangen, Germany

AUD-33.5: IMPROVING CHARACTER ERROR RATE IS NOT EQUAL TO HAVING CLEAN SPEECH: SPEECH ENHANCEMENT FOR ASR SYSTEMS WITH BLACK-BOX ACOUSTIC MODELS 991

Ryosuke Sawata, Yosuke Kashiwagi, Shusuke Takahashi, Sony Group Corporation, Japan

AUD-33.6: EFFECT OF NOISE SUPPRESSION LOSSES ON SPEECH DISTORTION AND ASR PERFORMANCE 996

Sebastian Braun, Hannes Gamper, Microsoft, United States of America

AUD-34: DECLIPPING, LOUDNESS, AND TIME DOMAIN MODIFICATION

AUD-34.1: INCREASING LOUDNESS IN AUDIO SIGNALS: A PERCEPTUALLY MOTIVATED APPROACH TO PRESERVE AUDIO QUALITY 1001

Alix Jeannerot, Paolo Prandoni, École Polytechnique Fédérale de Lausanne (EPFL), Switzerland; Niels de Koeijer, Delft University of Technology (TU Delft), Netherlands; Pablo Martínez-Nuevo, Martin Bo Møller, Jakob Dyreby, Bang & Olufsen, Denmark

AUD-34.2: AUDIO PEAK REDUCTION USING A SYNCED ALLPASS FILTER 1006

Sebastian J. Schlecht, Leonardo Fierro, Vesa Välimäki, Aalto University, Finland; Juha Backman, AAC Technologies Solutions, Finland

AUD-34.3: APPLADE: ADJUSTABLE PLUG-AND-PLAY AUDIO DECLIPPER COMBINING DNN WITH SPARSE OPTIMIZATION 1011

Tomoro Tanaka, Kohei Yatabe, Yasuhiro Oikawa, Waseda University, Japan; Masahiro Yasuda, NTT Corporation, Japan

: DETECTION AND CLASSIFICATION OF ACOUSTIC SCENES AND EVENTS VII: EVALUATION

.1: MAXIMIZING AUDIO EVENT DETECTION MODEL PERFORMANCE ON SMALL DATASETS THROUGH KNOWLEDGE TRANSFER, DATA AUGMENTATION, AND PRETRAINING: AN ABLATION STUDY 1016

Daniel Tompkins, Kshitiz Kumar, Jian Wu, Microsoft, United States of America

.2: THRESHOLD INDEPENDENT EVALUATION OF SOUND EVENT DETECTION SCORES 1021

Janek Ebbers, Reinhold Haeb-Umbach, Paderborn University, Germany; Romain Serizel, Université de Lorraine, France

.3: MULTIMODAL EVALUATION METHOD FOR SOUND EVENT DETECTION.....	1026
<i>Seyed Mohammad Reza Modaresi, Aomar Osmani, Sorbonne Paris Nord University, France; Mohammadreza Razzazi, Amirkabir University of Technology, Iran (Islamic Republic of); Abdelghani Chibani, Université Paris-Est Créteil, France</i>	

.4: A BENCHMARK OF STATE-OF-THE-ART SOUND EVENT DETECTION SYSTEMS	1031
EVALUATED ON SYNTHETIC SOUNDSCAPES	
<i>Francesca Ronchini, Romain Serizel, INRIA Grand Est, France</i>	

AUD-35: TOPICS IN CONVOLUTIONAL AND LATENT VARIABLE MODELING

AUD-35.1: ATTENTIVE MAX FEATURE MAP AND JOINT TRAINING FOR ACOUSTIC	1036
SCENE CLASSIFICATION	
<i>Hye-jin Shim, Ju-ho Kim, Ha-Jin Yu, University of Seoul, Korea, Republic of; Jee-weon Jung, Naver corporation, Korea, Republic of</i>	

AUD-35.2: A VARIATIONAL BAYESIAN APPROACH TO LEARNING LATENT VARIABLES	1041
FOR ACOUSTIC KNOWLEDGE TRANSFER	
<i>Hu Hu, Chao-Han Huck Yang, Chin-Hui Lee, Georgia Institute of Technology, United States of America; Sabato Marco Siniscalchi, University of Enna Kore, Italy</i>	

AUD-35.5: ORCA-PARTY: AN AUTOMATIC KILLER WHALE SOUND TYPE SEPARATION	1046
TOOLKIT USING DEEP LEARNING	
<i>Christian Bergler, Manuel Schmitt, Andreas Maier, Elmar Nöth, Friedrich-Alexander-University Erlangen-Nuremberg, Germany; Rachael Xi Cheng, Leibniz Institute for Zoo and Wildlife Research, Germany; Volker Barth, Anthro-Media, Germany</i>	

AUD-36: SPATIAL AUDIO II

AUD-36.1: SPARSITY-BASED SOUND FIELD SEPARATION IN THE SPHERICAL	1051
HARMONICS DOMAIN	
<i>Mirco Pezzoli, Fabio Antonacci, Augusto Sarti, Politecnico di Milano, Italy; Maximo Cobos, Universitat de Valencia, Spain</i>	

AUD-36.2: SPATIAL ACTIVE NOISE CONTROL BASED ON INDIVIDUAL KERNEL	1056
INTERPOLATION OF PRIMARY AND SECONDARY SOUND FIELDS	
<i>Kazuyuki Arikawa, Shoichi Koyama, Hiroshi Saruwatari, The University of Tokyo, Japan</i>	

AUD-36.3: TIME-DOMAIN ACOUSTIC CONTRAST CONTROL WITH A SPATIAL	1061
UNIFORMITY CONSTRAINT FOR PERSONAL AUDIO SYSTEMS	
<i>Sipei Zhao, Ian Burnett, University of Technology Sydney, Australia</i>	

AUD-36.4: GENERATION OF PERSONAL SOUND FIELDS IN REVERBERANT	1066
ENVIRONMENTS USING INTERFRAME CORRELATION	
<i>Liming Shi, Mads Graesboll Christensen, Aalborg University, Denmark; Guoli Ping, Xiaoxiang Shen, Huawei Technologies Co., Ltd, Denmark</i>	

AUD-36.5: VARIABLE SPAN TRADE-OFF FILTER FOR SOUND ZONE CONTROL	1071
WITH KERNEL INTERPOLATION WEIGHTING	
<i>Jesper Brunnström, Marc Moonen, KU Leuven, Belgium; Shoichi Koyama, The University of Tokyo, Japan</i>	

AUD-36.6: TIME DOMAIN RADIAL FILTER DESIGN FOR SPHERICAL WAVES.....	1076
<i>Nara Hahn, Frank Schultz, Sascha Spors, University of Rostock, Germany</i>	

BIO-1: MEDICAL IMAGE SEGMENTATION I

BIO-1.1: FEATURE SPACE MESSAGE PASSING NETWORK FOR MEDICAL IMAGE SEMANTIC SEGMENTATION 1081

Junxiao Sun, Ke Zhang, Shuyi Niu, Yan Zhang, Youyong Kong, Southeast University, China

BIO-1.2: CROSS-DOMAIN FEW-SHOT LEARNING FOR RARE-DISEASE SKIN LESION SEGMENTATION 1086

Yixin Wang, Institute of Computing Technology, Chinese Academy of Sciences, China; Zhe Xu, Department of Biomedical Engineering, The Chinese University of Hong Kong, China; Jiang Tian, Zhongchao Shi, AI Lab, Lenovo Research, China; Jie Luo, Brigham and Women's Hospital, Harvard Medical School, China; Yang Zhang, Zhiqiang He, Lenovo Ltd, China; Jianping Fan, UNCC; AI Lab, Lenovo Research, China

BIO-1.3: ADAPTIVE PSEUDO LABELING FOR SOURCE-FREE DOMAIN ADAPTATION IN MEDICAL IMAGE SEGMENTATION 1091

Chen Li, Wei Chen, Xin Luo, Yulin He, Yusong Tan, National University of Defense Technology, China

BIO-1.4: OBJECT DETECTION AND TRACKING IN ULTRASOUND SCANS USING AN OPTICAL FLOW AND SEMANTIC SEGMENTATION FRAMEWORK BASED ON CONVOLUTIONAL NEURAL NETWORKS 1096

Abdullah Al-Battal, Imanuel Lerman, Truong Nguyen, University of California - San Diego, United States of America

BIO-1.5: HEURISTIC DROPOUT: AN EFFICIENT REGULARIZATION METHOD FOR MEDICAL IMAGE SEGMENTATION MODELS 1101

Dachuan Shi, Ruiyang Liu, Linmi Tao, Chun Yuan, Tsinghua University, China

BIO-1.6: SUPERRESOLUTION AND SEGMENTATION OF OCT SCANS USING MULTI-STAGE ADVERSARIAL GUIDED ATTENTION TRAINING 1106

Paria Jekhouni, Omid Dehzangi, Nasser M. Nasrabadi, West Virginia University, United States of America; Annahita Amireskandari, Ali Dabouei, Ali Rezai, West Virginia University, United States of America

BIO-2: MACHINE LEARNING FOR PHYSIOLOGICAL SIGNALS I

BIO-2.1: HEART RATE AND OXYGEN SATURATION ESTIMATION FROM FACIAL VIDEO WITH MULTIMODAL PHYSIOLOGICAL DATA GENERATION 1111

Yusuke Akamatsu, Yoshifumi Onishi, Hitoshi Imaoka, NEC Corporation, Japan

BIO-2.2: EMGSE: ACOUSTIC/EMG FUSION FOR MULTIMODAL SPEECH ENHANCEMENT 1116

Kuan-Chen Wang, National Taiwan University, Taiwan; Kai-Chun Liu, Hsin-Min Wang, Yu Tsao, Academia Sinica, Taiwan

BIO-2.3: A DILATED RESIDUAL VISION TRANSFORMER FOR ATRIAL FIBRILLATION DETECTION FROM STACKED TIME-FREQUENCY ECG REPRESENTATIONS 1121

Sawon Pratiher, Apoorva Srivastava, Yedla Bindu Priyatha, Nirmalya Ghosh, Amit Patra, Indian Institute of Technology Kharagpur, India

BIO-2.4: CONTRASTIVE HEARTBEATS: CONTRASTIVE LEARNING FOR SELF-SUPERVISED ECG REPRESENTATION AND PHENOTYPING 1126

Crystal T. Wei, Ming-En Hsieh, Chien-Liang Liu, Vincent S. Tseng, National Yang Ming Chiao Tung University, Taiwan, Province of China

BIO-2.5: UBIQUITOUS PHYSIOLOGICAL PREDICTION OF SUD PATIENTS' WELLNESS STATE USING MEMORY-BASED CONVOLUTIONAL MODELS 1131

Omid Dehzangi, Paria Jekhouni, Victor Finomore, Nasser M. Nasrabadi, Ali Rezai, West Virginia University, United States of America; Jad Ramadan, West Virginia University, United States of America

BIO-2.6: JOINT HYPOGLYCEMIA PREDICTION AND GLUCOSE FORECASTING VIA DEEP MULTI-TASK LEARNING	1136
<i>Mu Yang, University of Texas at Dallas, United States of America; Darpit Dave, Madhav Erraguntla, Gerard Cote, Ricardo Gutierrez-Osuna, Texas A&M University, United States of America</i>	
BIO-3: DEEP LEARNING FOR PHYSIOLOGICAL SIGNAL ANALYSIS I	
BIO-3.1: SEGNET-BASED DEEP REPRESENTATION LEARNING FOR DYSPHAGIA CLASSIFICATION	1141
<i>Siddharth Subramani, Achuth Rao M V, Anwesha Roy, Prasanta Ghosh, Indian Institute of Science, India; Prasanna Suresh Hegde, Health Care Global Enterprises Ltd, India</i>	
BIO-3.2: ROBUST COLLABORATIVE LEARNING FOR SEQUENCE MODELLING	1146
<i>Francois Buet-Golfouse, University College London, United Kingdom of Great Britain and Northern Ireland; Hans Roggeman, Islam Utyagulov, Independent, United Kingdom of Great Britain and Northern Ireland</i>	
BIO-3.3: A SELF-SUPERVISED PRE-TRAINING FRAMEWORK FOR VISION-BASED SEIZURE CLASSIFICATION	1151
<i>Jen-Cheng Hou, Monique Thonnat, University of Côte d'Azur, INRIA, France; Aileen McGonigal, Fabrice Bartolomei, Aix-Marseille University, France</i>	
BIO-3.4: DESIGN OF REAL-TIME SYSTEM BASED ON MACHINE LEARNING FOR SNORING AND OSA DETECTION	1156
<i>Huaiwen Luo, Lu Zhang, Lianyu Zhou, Xu Lin, Zehuai Zhang, Mingjiang Wang, Harbin Institute of Technology, Shenzhen, China</i>	
BIO-3.5: PARAMETRIC MODELING OF HUMAN WRIST FOR BIOIMPEDANCE-BASED PHYSIOLOGICAL SENSING	1161
<i>Kaan Sel, Noah Huerta, Roozbeh Jafari, Texas A&M University, United States of America; Michael S. Sacks, The University of Texas at Austin, United States of America</i>	
BIO-3.6: PRELIMINARY RESULTS ON THE GENERATION OF ARTIFICIAL HANDWRITING DATA USING A DECOMPOSITION-RECOMBINATION STRATEGY	1166
<i>Jose Fernando Adrán Otero, Oscar Solans Caballer, Open University of Catalonia, Spain; Pere Marti-Puig, Jordi Solé-Casals, University of Vic - Central University of Catalonia, Spain; Zhe Sun, RIKEN, Japan; Toshihisa Tanaka, Tokyo University of Agriculture and Technology, Japan</i>	
BIO-4: MACHINE LEARNING FOR PHYSIOLOGICAL SIGNALS II	
BIO-4.1: A STYLE TRANSFER MAPPING AND FINE-TUNING SUBJECT TRANSFER FRAMEWORK USING CONVOLUTIONAL NEURAL NETWORKS FOR SURFACE ELECTROMYOGRAM PATTERN RECOGNITION	1171
<i>Suguru Kanoga, Takayuki Hoshino, Mitsunori Tada, National Institute of Advanced Industrial Science and Technology (AIST), Japan</i>	
BIO-4.2: FEATURE-BASED SENSING MATRIX DESIGN FOR ANALOG TO INFORMATION CONVERTERS	1176
<i>Chencheng Guo, Hui Qian, Baoling Hong, Fuzhou University, China</i>	
BIO-4.3: ALSNET: A DILATED 1-D CNN FOR IDENTIFYING ALS FROM RAW EMG SIGNAL	1181
<i>K. M. Naimul Hassan, Md. Shamiul Alam Hridoy, Naima Tasnim, Atia Faria Chowdhury, Tanvir Alam Roni, Sheikh Tabrez, Arik Subhana, Celia Shahnaz, Bangladesh University of Engineering and Technology, Bangladesh</i>	

BIO-4.4: JOINT MODEL ORDER ESTIMATION FOR MULTIPLE TENSORS WITH A COUPLED MODE AND APPLICATIONS TO THE JOINT DECOMPOSITION OF EEG, MEG MAGNETOMETER, AND GRADIOMETER TENSORS	1186
<i>Bilal Ahmad, Liana Khamidullina, Alla Manina, Jens Haueisen, Martin Haardt, Ilmenau University of Technology, Germany; Alexey Korobkov, Kazan National Research Technical University n.a. A.N. Tupolev-KAI, Russian Federation</i>	
BIO-4.5: AN EXPERIMENTAL STUDY ON TRANSFERRING DATA-DRIVEN IMAGE COMPRESSIVE SENSING TO BIOELECTRIC SIGNAL	1191
<i>Zhikang Zhang, Jonathan Zhao, Fengbo Ren, Arizona State University, United States of America</i>	
BIO-4.6: HAND GESTURE RECOGNITION USING TEMPORAL CONVOLUTIONS AND ATTENTION MECHANISM	1196
<i>Elahe Rahimian, Soheil Zabihi, Amir Asif, Arash Mohammadi, Concordia University, Canada; Dario Farina, Imperial College London, United Kingdom of Great Britain and Northern Ireland; Farokh Atashzar, New York University, United States of America</i>	
BIO-5: DEEP LEARNING FOR MEDICAL IMAGE ANALYSIS	
BIO-5.1: COMBINING MULTIPLE STYLE TRANSFER NETWORKS AND TRANSFER LEARNING FOR LGE-CMR SEGMENTATION	1201
<i>Bo Fang, Junxin Chen, Northeastern University, China; Wei Wang, Sun Yat-sen University, China; Yicong Zhou, University of Macau, China</i>	
BIO-5.2: MULTI-DOMAIN UNPAIRED ULTRASOUND IMAGE ARTIFACT REMOVAL USING A SINGLE CONVOLUTIONAL NEURAL NETWORK	1206
<i>Jaeyoung Huh, Shujaat Khan, Jong Chul Ye, Korea Advanced Institute of Science and Technology, Korea, Republic of</i>	
BIO-5.3: IMPROVING ULTRASOUND IMAGE CLASSIFICATION WITH LOCAL TEXTURE QUANTISATION	1211
<i>Xiao Li, Huizhi Liang, University of Reading, United Kingdom of Great Britain and Northern Ireland; Sidhartha Nagala, Jane Chen, Royal Berkshire Hospital, United Kingdom of Great Britain and Northern Ireland</i>	
BIO-5.4: ACCELERATED INTRAVASCULAR ULTRASOUND IMAGING USING DEEP REINFORCEMENT LEARNING	1216
<i>Tristan Stevens, Nishith Chennakeshava, Ruud van Sloun, Eindhoven University of Technology, Netherlands; Frederik de Bruijn, Martin Pekař, Philips Research Eindhoven, Netherlands</i>	
BIO-5.5: DEEP PROXIMAL UNFOLDING FOR IMAGE RECOVERY FROM UNDER-SAMPLED CHANNEL DATA IN INTRAVASCULAR ULTRASOUND	1221
<i>Nishith Chennakeshava, Tristan S.W. Stevens, Massimo Mischi, Eindhoven University of Technology, Netherlands; Frederik J. de Bruijn, Andrew Hancock, Martin Pekař, Philips Research, Netherlands; Yonina C. Eldar, The Weizmann Institute of Science, Netherlands; Ruud J.G. van Sloun, Eindhoven University of Technology, Philips Research, Netherlands</i>	
BIO-5.6: MULTIVIEW LONG-SHORT SPATIAL CONTRASTIVE LEARNING FOR 3D MEDICAL IMAGE ANALYSIS	1226
<i>Gongpeng Cao, Yiping Wang, Manli Zhang, Jing Zhang, Guixia Kang, Beijing University of Posts and Telecommunications, China; Xin Xu, General Hospital of People's Liberation Army, China</i>	
BIO-6: MACHINE LEARNING AND SIGNAL PROCESSING FOR EEG	
BIO-6.1: COMPOSING GRAPHICAL MODELS WITH GENERATIVE ADVERSARIAL NETWORKS FOR EEG SIGNAL MODELING	1231
<i>Khuong Vo, Manoj Vishwanath, Ramesh Srinivasan, Nikil Dutt, Hung Cao, University of California, Irvine, United States of America</i>	

BIO-6.2: DOMAIN-INVARIANT REPRESENTATION LEARNING FROM EEG WITH PRIVATE ENCODERS	1236
<i>David Bethge, Porsche, LMU, Germany; Philipp Hallgarten, Porsche, KIT, Germany; Tobias Grosse-Puppendahl, Mohamed Kari, Porsche, Germany; Ralf Mikut, KIT, Germany; Albrecht Schmidt, LMU, Germany; Ozan Özdenizci, TU Graz, Austria</i>	
BIO-6.3: HOLISTIC SEMI-SUPERVISED APPROACHES FOR EEG REPRESENTATION LEARNING	1241
<i>Guangyi Zhang, Ali Etemad, Queen's University, Canada</i>	
BIO-6.4: MUSIC IDENTIFICATION USING BRAIN RESPONSES TO INITIAL SNIPPETS	1246
<i>Pankaj Pandey, Krishna. P. Miyapuram, IIT Gandhinagar, India; Gulshan Sharma, IIT Ropar, India; Ramanathan Subramanian, University of Canberra, Australia; Derek Lomas, TU Delft, Netherlands</i>	
BIO-6.5: MULTI-LEVEL SPATIAL-TEMPORAL ADAPTATION NETWORK FOR MOTOR IMAGERY CLASSIFICATION	1251
<i>Wei Xu, Jing Wang, Ziyu Jia, Zhiqing Hong, Yunze Li, Youfang Lin, Beijing Jiaotong University, China</i>	
BIO-6.6: LEARNING SUBJECT-INVARIANT REPRESENTATIONS FROM SPEECH-EVOKED EEG USING VARIATIONAL AUTOENCODERS	1256
<i>Lies Bollens, Tom Francart, Hugo Van hamme, Katholieke Universiteit Leuven, Belgium</i>	
BIO-7: MACHINE LEARNING FOR MEDICAL IMAGING	
BIO-7.1: UNSUPERVISED HIERARCHICAL TRANSLATION-BASED MODEL FOR MULTI-MODAL MEDICAL IMAGE REGISTRATION	1261
<i>Xinru Dai, Tai Ma, Haibin Cai, Ying Wen, East China Normal University, China</i>	
BIO-7.2: FAZ-BV: A DIABETIC MACULAR ISCHEMIA GRADING FRAMEWORK COMBINING FAZ ATTENTION NETWORK AND BLOOD VESSEL ENHANCEMENT FILTERS	1266
<i>Zailiang Chen, Hailei Lan, Yuchen Xiong, Hailan Shen, School of Computer Science and Engineering, Central South University, Changsha, China, China; Yongan Meng, Jing Luo, The Second Xiangya Hospital, Central South University, Changsha, China, China</i>	
BIO-7.3: FRACTURE DETECTION AND LOCALIZATION IN CHEST X-RAYS USING SEMI-SUPERVISED LEARNING WITH DYNAMIC SHARPENING	1271
<i>Lijuan Lu, Tsinghua University, China; Shun Miao, PAII Inc., United States of America; Ling Ye, Ping An Technology, China</i>	
BIO-7.4: HISTOKT: CROSS KNOWLEDGE TRANSFER IN COMPUTATIONAL PATHOLOGY	1276
<i>Ryan Zhang, Jiadai Zhu, Stephen Yang, Konstantinos Plataniotis, University of Toronto, Canada; Mahdi Hosseini, University of New Brunswick, Canada; Angelo Genovese, Università degli Studi di Milano, Italy; Lina Chen, Sunnybrook Health Sciences Centre, Canada; Corwyn Rowsell, St. Michaels Hospital, Canada; Savvas Damaskinos, Huron Digital Pathology, Canada; Sonal Varma, Kingston Health Sciences Center, Canada</i>	
BIO-7.5: UNSUPERVISED DEEP LEARNING NETWORK FOR DEFORMABLE FUNDUS IMAGE REGISTRATION	1281
<i>Giovana Benvenuto, Marilaine Colnago, Wallace Casaca, São Paulo State University, Brazil</i>	
BIO-7.6: A MINIMALLY SUPERVISED APPROACH FOR MEDICAL IMAGE QUALITY ASSESSMENT IN DOMAIN SHIFT SETTINGS	1286
<i>Huijuan Yang, Feri Guretno, Ivan Ho Mien, Chuan Sheng Foo, Pavitra Krishnaswamy, Institute for Infocomm Research, Singapore; Aaron S. Coyner, J. Peter Campbell, Susan Ostmo, Oregon Health & Science University, Singapore; Michael F. Chiang, National Institutes of Health, United States of America</i>	

BIO-8: BRAIN-COMPUTER INTERFACES

BIO-8.1: A CHANNEL ATTENTION BASED MLP-MIXER NETWORK FOR MOTOR IMAGERY DECODING WITH EEG 1291

Yanbin He, Zhiyang Lu, Jun Wang, Jun Shi, Shanghai University, China

BIO-8.2: TOWARDS CLOSED-LOOP SPEECH SYNTHESIS FROM STEREOTACTIC EEG: A UNIT SELECTION APPROACH 1296

Miguel Angrick, Lorenz Diener, Darius Ivucic, Gabriel Ivucic, Tanja Schultz, University of Bremen, Germany; Maarten Ottenhoff, Sophocles Goulis, Pieter L. Kubben, Christian Herff, Maastricht University, Netherlands; Albert J. Colon, Louis Wagner, Epilepsy Center Kempenhaeghe, Netherlands; Dean J. Krusienski, Virginia Commonwealth University, United States of America

BIO-8.3: ENHANCING CONTEXTUAL ENCODING WITH STAGE-CONFUSION AND STAGE-TRANSITION ESTIMATION FOR EEG-BASED SLEEP STAGING 1301

Jaeun Phyoo, Wonjun Ko, Eunjin Jeon, Heung-Il Suk, Korea University, Korea, Republic of

BIO-8.4: IMPROVING BCI-BASED COLOR VISION ASSESSMENT USING GAUSSIAN PROCESS REGRESSION 1306

Hadi Habibzadeh, Daphney-Stavroula Zois, University at Albany, State University of New York, United States of America; Kevin J. Long, James J. S. Norton, US Department of Veterans Affairs, United States of America; Ally E. Atkins, Union College, United States of America

BIO-8.5: TRANSFORMER-BASED ESTIMATION OF SPOKEN SENTENCES USING ELECTROCORTICOGRAPHY 1311

Shuji Komeiji, Kai Shigemi, Toshihisa Tanaka, Tokyo University of Agriculture and Technology, Japan; Takumi Mitsuhashi, Yasushi Iimura, Hiroharu Suzuki, Hidenori Sugano, Juntendo University, Japan; Koichi Shinoda, Tokyo Institute of Technology, Japan

BIO-9: DEEP LEARNING FOR PHYSIOLOGICAL SIGNAL ANALYSIS II

BIO-9.1: BOOST ENSEMBLE LEARNING FOR CLASSIFICATION OF CTG SIGNALS 1316

Marzieh Ajirak, Petar Djuric, Stony Brook University, United States of America; Cassandra Heiselman, Gerald Quirk, Stony Brook Medicine, United States of America

BIO-9.2: MULTI-VIEW LEARNING BASED ON NON-REDUNDANT FUSION FOR ICU PATIENT MORTALITY PREDICTION 1321

Yifan Wang, Ying Lan, Peking University, China

BIO-9.3: IMPROVING PHASE-RECTIFIED SIGNAL AVERAGING FOR FETAL HEART RATE ANALYSIS 1326

Tong Chen, Guanchao Feng, Cassandra Heiselman, J. Gerald Quirk, Petar M. Djuric, Stony Brook University, United States of America

BIO-9.4: UNSUPERVISED CLUSTERING AND ANALYSIS OF CONTRACTION-DEPENDENT FETAL HEART RATE SEGMENTS 1331

Liu Yang, Petar M. Djurić, Stony Brook University, United States of America; Cassandra Heiselman, J. Gerald Quirk, Stony Brook University Hospital, United States of America

BIO-9.5: A METHOD FOR DETECTING CORONARY ARTERY DISEASE USING NOISY ULTRASHORT ELECTROCARDIOGRAM RECORDINGS 1336

Orestis Apostolou, Vasileios Charisis, Georgios Apostolidis, Aristotle University of Thessaloniki, Greece; Leontios J. Hadjileontiadis, Aristotle University of Thessaloniki, Khalifa University, Greece

BIO-9.6: MULTI-TASK GAUSSIAN PROCESS REGRESSION FOR THE DETECTION OF SLEEP CYCLES IN PREMATURE INFANTS	1341
<i>Nele Sophie Brügge, University of Lübeck, Germany; Jan Grasshoff, Fraunhofer Research Institution for Individualized and Cell-Based Medical Engineering, Germany; Arne Weigenand, Drägerwerk AG & Co. KGaA, Germany; Philipp Rostalski, University of Lübeck and Fraunhofer Research Institution for Individualized and Cell-Based Medical Engineering, Germany</i>	
BIO-10: DETECTION, ESTIMATION AND CLASSIFICATION FOR MEDICAL IMAGES	
BIO-10.1: FAST LOW RANK COLUMN-WISE COMPRESSIVE SENSING FOR ACCELERATED DYNAMIC MRI	1346
<i>Silpa Babu, Seyedehsara Nayer, Namrata Vaswani, Iowa State University, United States of America; Sajan Lingala, University of Iowa, United States of America</i>	
BIO-10.2: MRI RECOVERY WITH A SELF-CALIBRATED DENOISER	1351
<i>Sizhuo Liu, Philip Schniter, Rizwan Ahmad, Ohio State University, United States of America</i>	
BIO-10.3: 3D CROSS-SCALE FEATURE TRANSFORMER NETWORK FOR BRAIN MR IMAGE SUPER-RESOLUTION	1356
<i>Wanqi Zhang, Lulu Wang, Wei Chen, Zhongshi He, Chongqing University, China; Yuanyuan Jia, Jinglong Du, Chongqing Medical University, China</i>	
BIO-10.4: DATA EFFICIENT SUPPORT VECTOR MACHINE TRAINING USING THE MINIMUM DESCRIPTION LENGTH PRINCIPLE	1361
<i>Harsh Singh, International Institute of Information Technology, India; Ognjen Arandjelovic, University of St Andrews, United Kingdom of Great Britain and Northern Ireland</i>	
BIO-10.5: MULTIPLE INSTANCE LEARNING WITH TASK-SPECIFIC MULTI-LEVEL FEATURES FOR WEAKLY ANNOTATED HISTOPATHOLOGICAL IMAGE CLASSIFICATION	1366
<i>Yuanpin Zhou, Yao Lu, Sun Yat-sen University, China</i>	
BIO-11: MEDICAL IMAGE ANALYSIS FOR DIAGNOSIS	
BIO-11.1: SELF-KNOWLEDGE DISTILLATION BASED SELF-SUPERVISED LEARNING FOR COVID-19 DETECTION FROM CHEST X-RAY IMAGES	1371
<i>Guang Li, Ren Togo, Takahiro Ogawa, Miki Haseyama, Hokkaido University, Japan</i>	
BIO-11.2: PIXEL-LEVEL AND AFFINITY-LEVEL KNOWLEDGE DISTILLATION FOR UNSUPERVISED SEGMENTATION OF COVID-19 LESIONS	1376
<i>Rui Xu, Yufeng Wang, Xinchun Ye, Pengcheng Wu, Dalian University of Technology, China; Yen-Wei Chen, Ritsumeikan University, Japan; Fangyi Xu, Wenchao Zhu, Chao Chen, Yong Zhou, Hongjie Hu, Zhejiang University School of Medicine, China; Xiaofeng Qu, the Second Hospital of Dalian Medical University, China; Shoji Kido, Noriyuki Tomiyama, Osaka University, Japan</i>	
BIO-11.3: DATA SHAPLEY VALUE FOR HANDLING NOISY LABELS: AN APPLICATION IN SCREENING COVID-19 PNEUMONIA FROM CHEST CT SCANS	1381
<i>Nastaran Enshaei, Arash Mohammadi, Farnoosh Naderkhani, Concordia University, Canada; Moezedin Javad Rafiee, McGill University, Canada</i>	
BIO-11.4: ACCURATE MULTISCALE SELECTIVE FUSION OF CT AND VIDEO IMAGES FOR REAL-TIME ENDOSCOPIC CAMERA 3D TRACKING IN ROBOTIC SURGERY	1386
<i>Xiongbiao Luo, Xiamen University, China</i>	
BIO-11.5: LEARNING DEEP PATHOLOGICAL FEATURES FOR WSI-LEVEL CERVICAL CANCER GRADING	1391
<i>Ruixiang Geng, Qing Liu, Shuo Feng, Yixiong Liang, Central South University, China</i>	

BIO-11.6: SELECTIVE SCALE CASCADE ATTENTION NETWORK FOR BREAST CANCER HISTOPATHOLOGY IMAGE CLASSIFICATION	1396
<i>Bowen Xu, Wenqiang Zhang, Fudan University, China</i>	

BIO-12: NEURAL SIGNAL PROCESSING

BIO-12.1: FREQUENCY-SPECIFIC NON-LINEAR GRANGER CAUSALITY IN A NETWORK OF BRAIN SIGNALS	1401
--	-------------

Archishman Biswas, Hernando Ombao, King Abdullah University of Science and Technology, Saudi Arabia

BIO-12.2: EPILEPTIC SPIKE DETECTION BY RECURRENT NEURAL NETWORKS WITH SELF-ATTENTION MECHANISM	1406
---	-------------

Kosuke Fukumori, Toshihisa Tanaka, Tokyo University of Agriculture and Technology, Japan; Noboru Yoshida, Juntendo University Nerima Hospital, Japan; Hidenori Sugano, Madoka Nakajima, Juntendo University Hospital, Japan

BIO-12.3: TOPOLOGICAL CORRELATION OF BRAIN SIGNALS	1411
---	-------------

Jian Yin, Nanjing Medical University, China; Yuan Wang, University of South Carolina, United States of America

BIO-12.4: ONLINE DETECTION OF SCALP-INVISIBLE MESIAL-TEMPORAL BRAIN INTERICTAL EPILEPTIFORM DISCHARGES FROM EEG	1416
--	-------------

Bahman Abdi-Sargezeh, Saeid Sanei, Nottingham Trent University, United Kingdom of Great Britain and Northern Ireland; Antonio Valentin, King's College London, United Kingdom of Great Britain and Northern Ireland; Gonzalo Alarcon, Hamad General Hospital, Qatar

BIO-12.5: LEVERAGING SPARSE CODING FOR EEG BASED EMOTION RECOGNITION IN SHOOTING	1421
---	-------------

Yulu Wang, Yiwen Sun, Changshui Zhang, Tsinghua University, China; Fang Lei, DataCanvas Technology Co., Ltd., China

BIO-12.6: A NOVEL UNSUPERVISED AUTOENCODER-BASED HFOS DETECTOR IN INTRACRANIAL EEG SIGNALS	1426
---	-------------

Weilai Li, Lanfeng Zhong, Weixi Xiang, Tongzhou Kang, Dakun Lai, University of Electronic Science and Technology of China, China

BIO-13: MEDICAL IMAGE SEGMENTATION II

BIO-13.1: A NOVEL CONVOLUTIONAL NEURAL NETWORK BASED ON ADAPTIVE MULTI-SCALE AGGREGATION AND BOUNDARY-AWARE FOR LATERAL VENTRICLE SEGMENTATION ON MR IMAGES	1431
--	-------------

Fei Ye, Kai Hu, Xiangtan University, China; Zhiqiang Wang, Sheng Zhu, Affiliated Hospital of Xiangnan University, China; Xuanya Li, Baidu Inc., China

BIO-13.2: MULTISCALE ATTENTION AGGREGATION NETWORK FOR 2D VESSEL SEGMENTATION	1436
--	-------------

Wentao Liu, Huihua Yang, Weijin Xu, School of Artificial Intelligence, Beijing University of Posts and Telecommunications, China; Tong Tian, the State Key Laboratory of Structural Analysis for Industrial Equipment, School of Aeronautics and Astronautics, Dalian University of Technology, China; Xipeng Pan, the School of Computer Science and Information Security, Guilin University of Electronic Technology, Guilin 541004, China, China

BIO-13.3: TCRNET: MAKE TRANSFORMER, CNN AND RNN COMPLEMENT EACH OTHER	1441
--	-------------

Xinxin Shan, Tai Ma, Anqi Gu, Haibin Cai, Ying Wen, East China Normal University, China

BIO-13.4: DOUBLE NOISE MEAN TEACHER SELF-ENSEMBLING MODEL FOR SEMI-SUPERVISED TUMOR SEGMENTATION	1446
---	-------------

Ke Zheng, Junhai Xu, Jianguo Wei, Tianjin University, China

BIO-13.5: RETHINKING COMPUTER-AIDED PELVIS SEGMENTATION 1451
Siming Yuan, Qing Liu, Shenghui Liao, Fuchang Han, Central South University, China; Haitao Wei, Central South E3D Digital Medical and Virtual Reality Research Center, China; Yingqi Zhang, Tongji University, China

BIO-13.6: VISION TRANSFORMER-BASED RETINA VESSEL SEGMENTATION WITH DEEP ADAPTIVE GAMMA CORRECTION 1456
Hyunwoo Yu, Jae-hun Shim, Jaeho Kwak, Jou won Song, Suk-Ju Kang, Sogang University, Korea, Republic of

BIO-14: NEUROIMAGING DATA ANALYSIS

BIO-14.1: SPECTRAL PERMUTATION TEST ON PERSISTENCE DIAGRAMS 1461
Yuan Wang, Julius Fridriksson, University of South Carolina, United States of America; Moo Chung, University of Wisconsin - Madison, United States of America

BIO-14.2: MULTI-TASK FMRI DATA FUSION USING IVA AND PARAFAC2 1466
Isabell Lehmann, Tanuj Hasija, Peter J. Schreier, Paderborn University, Germany; Evrim Acar, Simula Research Laboratory, Norway; M.A.B.S. Akhonda, University of Maryland, Baltimore County, United States of America; Vince D. Calhoun, Tri-Institutional Georgia State University/Georgia Institute of Technology/Emory University Center, United States of America; Tülay Adalı, University of University of Maryland, Baltimore County, United States of America

BIO-14.3: INDEPENDENT VECTOR ANALYSIS BASED SUBGROUP IDENTIFICATION FROM MULTISUBJECT FMRI DATA 1471
Hanlu Yang, Mohammad Abu Baker Siddique Akhonda, Fatemeh Ghayyem, Qunfang Long, Tülay Adalı, University of Maryland Baltimore County, United States of America; Vince Calhoun, Georgia State University, Georgia Institute of Technology, and Emory University, United States of America

BIO-14.4: IMPROVING BRAIN DECODING METHODS AND EVALUATION 1476
Damián Pascual, Béni Egressy, Nicolas Affolter, Yiming Cai, Oliver Richter, Roger Wattenhofer, ETH Zurich, Switzerland

BIO-14.5: CMRI2SPEC: CINE MRI SEQUENCE TO SPECTROGRAM SYNTHESIS VIA A PAIRWISE HETEROGENEOUS TRANSLATOR 1481
Xiaofeng Liu, Jangwon Kim, Georges El Fakhri, Jonghye Woo, Harvard Medical School, United States of America; Fangxu Xing, MGH and Harvard Medical School, United States of America; Maureen Stone, UMD, United States of America; Jerry Prince, JHU, United States of America

BIO-14.6: SPATIO-TEMPORAL ATTENTION GRAPH CONVOLUTION NETWORK FOR FUNCTIONAL CONNECTOME CLASSIFICATION 1486
Wenhan Wang, Youyong Kong, Zhenghua Hou, Chunfeng Yang, Yonggui Yuan, Southeast University, China

CI-1: COMPUTATIONAL IMAGING METHODS AND MODELS

CI-1.1: BILEVEL LEARNING OF L1 REGULARIZERS WITH CLOSED-FORM GRADIENTS (BLORC) 1491
Avrajit Ghosh, Saiprasad Ravishankar, Michigan State University, United States of America; Michael Mccann, Los Alamos National Laboratory, United States of America

CI-1.2: MULTIBAND IMAGE FUSION WITH CONTROLLABLE ERROR GUARANTEES 1496
Unni V. S., Ruturaj Gavaskar, Kunal Chaudhury, Indian Institute of Science, India

CI-1.3: WEIGHTED GRAPH EMBEDDED LOW-RANK PROJECTION LEARNING FOR FEATURE EXTRACTION 1501
Zhuojie Huang, Shuping Zhao, Lunke Fei, Jigang Wu, Guangdong University of Technology, China

CI-1.4: ADMM-DAD NET: A DEEP UNFOLDING NETWORK FOR ANALYSIS	1506
COMPRESSED SENSING	
<i>Vasiliki Kouni, George C. Alexandropoulos, National and Kapodistrian University of Athens, Greece; Georgios Paraskevopoulos, National Technical University of Athens, Greece; Holger Rauhut, Rheinisch-Westfälische Technische Hochschule Aachen University, Germany</i>	
CI-1.5: HIGH-DIMENSIONAL SPARSE BAYESIAN LEARNING WITHOUT COVARIANCE	1511
MATRICES	
<i>Alexander Lin, Demba Ba, Harvard University, United States of America; Andrew Song, Massachusetts Institute of Technology, United States of America; Berkin Bilgic, Athinoula A. Martinos Center for Biomedical Imaging, United States of America</i>	
CI-1.6: A TRAINABLE BOUNDED DENOISER USING DOUBLE TIGHT FRAME	1516
NETWORK FOR SNAPSHOT COMPRESSIVE IMAGING	
<i>Baoshun Shi, Yuxin Wang, Qiusheng Lian, Yanshan University, China</i>	
CI-1.7: PROGRESSIVE IMAGE SUPER-RESOLUTION VIA NEURAL DIFFERENTIAL	1521
EQUATION	
<i>Seobin Park, Tae Hyun Kim, Hanyang University, Korea, Republic of</i>	
 CI-2: COMPUTATIONAL IMAGING FORMATION	
CI-2.1: HIGH-QUALITY SELF-SUPERVISED SNAPSHOT HYPERSPECTRAL IMAGING	1526
<i>Yuhui Quan, Xinran Qin, Mingqin Chen, Yan Huang, South China University of Technology, China</i>	
CI-2.2: ROBUST BAYESIAN RECONSTRUCTION OF MULTISPECTRAL	1531
SINGLE-PHOTON 3D LIDAR DATA WITH NON-UNIFORM BACKGROUND	
<i>Abderrahim Halimi, Jakeoung Koo, Gerald S. Buller, Stephen McLaughlin, Heriot-Watt University, United Kingdom of Great Britain and Northern Ireland; Robert Lamb, Leonardo MW Ltd., United Kingdom of Great Britain and Northern Ireland</i>	
CI-2.3: JOINT CALIBRATION AND MAPPING OF SATELLITE ALTIMETRY DATA USING	1536
TRAINABLE VARIATIONAL MODELS	
<i>Quentin Febvre, Ronan Fablet, IMT Atlantique, France; Julien Le Sommer, Université Grenoble Alpes, France; Clément Ubelmann, OceanNext, France</i>	
CI-2.4: 4D CONVOLUTIONAL NEURAL NETWORKS FOR MULTI-SPECTRAL AND	1541
MULTI-TEMPORAL REMOTE SENSING DATA CLASSIFICATION	
<i>Michalis Giannopoulos, Panagiotis Tsakalides, University of Crete, Greece; Grigorios Tsagkatakis, Foundation for Research and Technology Hellas, Greece</i>	
CI-2.5: A NEW DEEP LEARNING METHOD FOR MULTISPECTRAL IMAGE TIME	1546
SERIES COMPLETION USING HYPERSPECTRAL DATA	
<i>Cheick T. Cissé, Ahed Alboody, Matthieu Puigt, Gilles Roussel, Univ. Littoral Côte d'Opale, France; Vincent Vantrepotte, Cédric Jamet, Trung-Kien Tran, Univ. Littoral Côte d'Opale, CNRS, France</i>	
CI-2.6: IMAGE DENOISING WITH DEEP UNFOLDING AND NORMALIZING FLOWS	1551
<i>Xinyi Wei, Hans van Gorp, Lizeth Gonzalez Carabarin, Ruud van Sloun, Eindhoven University of Technology, Netherlands; Daniel Freedman, Verily Research, Israel; Yonina Eldar, Weizmann Institute of Science, Israel</i>	
CI-2.7: 3D TEXTURE SUPER RESOLUTION VIA THE RENDERING LOSS	1556
<i>Rohit Ranade, Yangwen Liang, Shuangquan Wang, Dongwoon Bai, Jungwon Lee, Samsung Semiconductor Inc., United States of America</i>	
 CI-3: COMPUTATIONAL PHOTOGRAPHY	
CI-3.1: BUNDLE ICP WITH VIRTUAL DEPTH FOR HAND-HELD 3D SCANNER	1561
<i>Changhun Sung, Byungdeok Kim, SAMSUNG ELECTRONICS, Korea, Republic of</i>	

CI-3.2: SKETCHED RT3D: HOW TO RECONSTRUCT BILLIONS OF PHOTONS PER SECOND	1566
<i>Julian Tachella, Michael P. Sheehan, Mike E. Davies, University of Edinburgh, United Kingdom of Great Britain and Northern Ireland</i>	
CI-3.3: A GENERIC METHOD TO ESTIMATE CAMERA EXTRINSIC PARAMETERS	1571
<i>Naveen Kuruba, Neel Badadare, Vikram Narayan, Satish Putta, Visteon, India</i>	
CI-3.4: PHOTON-LIMITED DEBLURRING USING ALGORITHM UNROLLING	1576
<i>Yash Sanghvi, Abhiram Gnanasambandan, Stanley Chan, Purdue University, United States of America</i>	
CI-3.5: NEX + : NOVEL VIEW SYNTHESIS WITH NEURAL REGULARISATION OVER MULTI-PLANE IMAGES	1581
<i>Wenpeng Xing, Jie Chen, Hong Kong Baptist University, Hong Kong</i>	
 CI-4: COMPUTATIONAL IMAGING SYSTEMS	
CI-4.1: COMPRESSIVE SCANNING TRANSMISSION ELECTRON MICROSCOPY	1586
<i>Daniel Nicholls, Alex Robinson, Jack Wells, Mounib Bahri, Nigel Browning, University of Liverpool, United Kingdom of Great Britain and Northern Ireland; Amirafshar Moshtaghpour, Rosalind Franklin Institute, United Kingdom of Great Britain and Northern Ireland; Angus Kirkland, University of Oxford, United Kingdom of Great Britain and Northern Ireland</i>	
CI-4.2: DEEP ITERATIVE PHASE RETRIEVAL FOR PTYCHOGRAPHY	1591
<i>Simon Welker, Tal Peer, Timo Gerkmann, Universität Hamburg, Germany; Henry Chapman, Deutsches Elektronen-Synchrotron, Germany</i>	
CI-4.3: COMPRESSIVE PHASE RETRIEVAL BASED ON SPARSE LATENT GENERATIVE PRIORS	1596
<i>Vinayak Killedar, Chandra Sekhar Seelamantula, Indian Institute of Science, India</i>	
CI-4.4: MODEL-BASED RECONSTRUCTION FOR COLLIMATED BEAM ULTRASOUND SYSTEMS	1601
<i>Abdulrahman Alanazi, Gregory Buzzard, Charles Bouman, Purdue University-Main Campus, United States of America; Singanallur Venkatakrishnan, Hector Santos-Villalobos, Oak Ridge National Laboratory, United States of America</i>	
CI-4.5: LEARNED ACOUSTIC RECONSTRUCTION USING SYNTHETIC APERTURE FOCUSING	1606
<i>Tim Straubinger, Robert Xiao, Helge Rhodin, University of British Columbia, Canada</i>	
CI-4.6: SDETR: ATTENTION-GUIDED SALIENT OBJECT DETECTION WITH TRANSFORMER	1611
<i>Guanze Liu, Bo Xu, Han Huang, Yandong Guo, OPPO Research Institute, China; Cheng Lu, Xmotors, China</i>	
 IVMSP-1: IMAGE AND VIDEO CODING I	
IVMSP-1.1: EVALUATION OF VIDEO CODING FOR MACHINES WITHOUT GROUND TRUTH	1616
<i>Kristian Fischer, André Kaup, Friedrich-Alexander-University Erlangen-Nürnberg, Germany; Markus Hofbauer, Christopher Kuhn, Eckehard Steinbach, Technical University of Munich, Germany</i>	
IVMSP-1.2: RAW PLENOPTIC VIDEO CODING UNDER HEXAGONAL LATTICE RESOLUTION OF MOTION VECTORS	1621
<i>Thuc Nguyen Huu, Vinh Duong Van, Jonghoon Yim, Byeungwoo Jeon, Sungkyunkwan University, Korea, Republic of</i>	
IVMSP-1.3: COMPARISON OF BOUNDARY ARTIFACT REMOVAL METHODS IN CODING OF GENERALIZED CUBEMAP PROJECTION USING VVC	1625
<i>Kianoush Jafari, Alireza Aminlou, Miska Hannuksela, Nokia Technologies, Finland</i>	

IVMSP-1.4: LOW-COMPLEXITY MULTI-MODEL CNN IN-LOOP FILTER FOR AVS3.....	1630
<i>Shen Wang, Yibing Fu, Chen Zhu, Li Song, Wenjun Zhang, Shanghai Jiao Tong University, China</i>	
IVMSP-1.5: UNIFIED MATRIX CODING FOR NN ORIGINATED MIP IN H.266/VVC	1635
<i>Junyan Huo, Yu Sun, Haixin Wang, Fuzheng Yang, Xidian University, China; Shuai Wan, Northwestern Polytechnical University, China; Ming Li, OPPO Mobile Telecommunications Corp., Ltd, China</i>	
IVMSP-1.6: FOV-BASED CODING OPTIMIZATION FOR 360-DEGREE VIRTUAL REALITY VIDEOS	1640
<i>Yuanyuan Xu, Taoyu Yang, Zengjie Tan, Haolun Lan, Hohai University, China</i>	
 IVMSP-2: MACHINE LEARNING FOR IMAGE PROCESSING I	
IVMSP-2.1: MULTI-HIERARCHY PROXY STRUCTURE FOR DEEP METRIC LEARNING	1645
<i>Jian Wang, Xinyue Li, Wei Song, Shanghai Ocean University, China; Zhichao Zhang, Shanghai Jiao Tong University, China; Weiqi Guo, East Sea Oceanographic Engineering Investigation, Design and Research Institute, China</i>	
IVMSP-2.2: EXPLOITING CAPTION DIVERSITY FOR UNSUPERVISED VIDEO SUMMARIZATION	1650
<i>Michail Kaseris, Ioannis Mademlis, Ioannis Pitas, Aristotle University of Thessaloniki, Greece</i>	
IVMSP-2.3: CLUSTERING AND SEPARATING SIMILARITIES FOR DEEP UNSUPERVISED HASHING	1655
<i>Wanqian Zhang, Dayan Wu, Bo Li, Weiping Wang, Institute of Information Engineering, Chinese Academy of Sciences, China; Chule Yang, National Innovation Institute of Defense Technology, Academy of Military Science, China</i>	
IVMSP-2.4: ENHANCING PROTOTYPICAL FEW-SHOT LEARNING BY LEVERAGING THE LOCAL-LEVEL STRATEGY	1660
<i>Junying Huang, Fan Chen, Keze Wang, Liang Lin, Dongyu Zhang, Sun Yat-sen University, China</i>	
IVMSP-2.5: BLIND UNMIXING USING A DOUBLE DEEP IMAGE PRIOR	1665
<i>Chao Zhou, Miguel Rodrigues, University College London, United Kingdom of Great Britain and Northern Ireland</i>	
IVMSP-2.6: A NEW FRAMEWORK FOR MULTIPLE DEEP CORRELATION FILTERS BASED OBJECT TRACKING	1670
<i>Yi Liu, Qiangqiang Wu, Hanzi Wang, Xiamen University, China; Yanjie Liang, Xiamen University, Peng Cheng Laboratory, China; Liming Zhang, University of Macau, China</i>	
 IVMSP-3: IMAGE PROCESSING	
IVMSP-3.1: ADAPTIVE ACTOR-CRITIC BILATERAL FILTER	1675
<i>Bo-Hao Chen, Hsiang-Yin Cheng, Yuan Ze University, Taiwan, Province of China; Jia-Li Yin, Fuzhou University, China</i>	
IVMSP-3.2: DOMAIN DECOMPOSITION ALGORITHMS FOR REAL-TIME HOMOGENEOUS DIFFUSION INPAINTING IN 4K	1680
<i>Niklas Kämper, Joachim Weickert, Saarland University, Germany</i>	
IVMSP-3.3: DEEP TEMPORAL INTERPOLATION OF RADAR-BASED PRECIPITATION.....	1685
<i>Michiaki Tatsubori, Takao Moriyama, Tatsuya Ishikawa, Paolo Fraccaro, Anne Jones, Blair Edwards, Julian Kuehnert, Sekou L. Remy, IBM Research, Japan</i>	
IVMSP-3.4: A NONLINEAR STEERABLE COMPLEX WAVELET DECOMPOSITION OF IMAGES	1690
<i>Zikai Sun, Thierry Blu, The Chinese University of Hong Kong, Hong Kong</i>	

IVMSP-3.5: KERNEL ESTIMATION NETWORK FOR BLIND SUPER-RESOLUTION 1695
Xiang Cao, Haibo Shen, Liangqi Zhang, Yihao Luo, Tianjiang Wang, Huazhong University of Science and Technology, China

IVMSP-3.6: TERAHERTZ IMAGE RESTORATION BENCHMARKING DATASET 1700
Yixiong Zhang, Zhipeng Su, School of Informatics Xiamen University (National Demonstrative Software School), Xiamen 361005, China, China; Feng Qi, Shenyang Institute of Automation, Chinese Academy of Sciences, Shenyang 110016, China, China; Jianyang Zhou, School of Electronic Science and Engineering, Xiamen University, Xiamen, China, China; Xiao-Ping Zhang, Department of Electrical, Computer and Biomedical Engineering, Ryerson University, Canada, Canada

IVMSP-4: IMAGE AND VIDEO UNDERSTANDING

IVMSP-4.1: BINARY DENSE PREDICTORS FOR HUMAN POSE ESTIMATION BASED ON DYNAMIC THRESHOLDS AND FILTERING 1705
Xingrun Xing, Yalong Jiang, Baochang Zhang, Wenrui Ding, Hongguang Li, Beihang University, China; Yangguang Li, SenseTime Group, China; Huan Peng, Huazhong University of Science and Technology, China

IVMSP-4.2: SELF-SUPERVISED LEARNING FOR SENTIMENT ANALYSIS VIA IMAGE-TEXT MATCHING 1710
Haidong Zhu, Zhaozheng Zheng, Ram Nevatia, University of Southern California, United States of America; Mohammad Soleymani, USC-ICT, United States of America

IVMSP-4.3: DOMAIN-AGNOSTIC META-LEARNING FOR CROSS-DOMAIN FEW-SHOT CLASSIFICATION 1715
Wei-Yu Lee, Jheng-Yu Wang, MOXA Inc., Taiwan; Yu-Chiang Wang, National Taiwan University, Taiwan

IVMSP-4.4: SEMANTIC ASSOCIATION NETWORK FOR VIDEO CORPUS MOMENT RETRIEVAL 1720
Dahyun Kim, Sunjae Yoon, Ji Woo Hong, Chang D Yoo, KAIST, Korea, Republic of

IVMSP-4.5: STATISTICAL, SPECTRAL AND GRAPH REPRESENTATIONS FOR VIDEO-BASED FACIAL EXPRESSION RECOGNITION IN CHILDREN 1725
Nida Itrat Abbasi, Siyang Song, Hatice Gunes, University of Cambridge, United Kingdom of Great Britain and Northern Ireland

IVMSP-4.6: DERIVING EXPLAINABLE DISCRIMINATIVE ATTRIBUTES USING CONFUSION ABOUT COUNTERFACTUAL CLASS 1730
Nakyeong Yang, Taegwan Kang, Kyomin Jung, Seoul National University, Korea, Republic of

IVMSP-5: IMAGE AND VIDEO SYNTHESIS

IVMSP-5.1: REALISTIC MONOCULAR-TO-3D VIRTUAL TRY-ON VIA MULTI-SCALE CHARACTERISTICS CAPTURE 1735
Chenghu Du, Feng Yu, Minghua Jiang, Yaxin Zhao, Xiong Wei, Tao Peng, Xinrong Hu, Wuhan Textile University, China

IVMSP-5.2: OPTIMIZING LATENT SPACE DIRECTIONS FOR GAN-BASED LOCAL IMAGE EDITING 1740
Ehsan Pajouheshgar, Tong Zhang, Sabine Süssstrunk, EPFL, Switzerland

IVMSP-5.3: TOWARDS USING CLOTHES STYLE TRANSFER FOR SCENARIO-AWARE PERSON VIDEO GENERATION 1745
Jingning Xu, School of Software Engineering, Tongji University, Shanghai, China, China; Benlai Tang, Wenyi Guo, Xiang Yin, Zejun Ma, ByteDance AI Lab, China; Mingjie Wang, School of Computer Science, University of Guelph, ON, Canada, Canada; Siyuan Bian, Department of Computer Science and Engineering, Shanghai Jiao Tong University, China, China

IVMSP-5.4: MULTI-DOMAIN UNSUPERVISED IMAGE-TO-IMAGE TRANSLATION WITH APPEARANCE ADAPTIVE CONVOLUTION	1750
<i>Somi Jeong, NAVER LABS, Korea, Republic of; Jiyoung Lee, NAVER AI Lab, Korea, Republic of; Kwanghoon Sohn, Yonsei University, Korea, Republic of</i>	
IVMSP-5.5: VR-FAM: VARIANCE-REDUCED ENCODER WITH NONLINEAR TRANSFORMATION FOR FACIAL ATTRIBUTE MANIPULATION	1755
<i>Yifan Yuan, Siteng Ma, Junping Zhang, Fudan University, China</i>	
IVMSP-5.6: WAVELET-BASED UNSUPERVISED LABEL-TO-IMAGE TRANSLATION	1760
<i>George Eskandar, Mohamed Abdelsamad, Karim Armanious, Shuai Zhang, Bin Yang, University of Stuttgart, Germany</i>	
IVMSP-6: VIDEO SUMMARIZATION AND ANOMALY DETECTION	
IVMSP-6.1: FAST GRAPH SAMPLING FOR SHORT VIDEO SUMMARIZATION USING GERSHGORIN DISC ALIGNMENT	1765
<i>Sadid Sahami, Chia-Wen Lin, National Tsing Hua University, Taiwan; Gene Cheung, York University, Canada</i>	
IVMSP-6.2: TOWARDS PRACTICAL AND EFFICIENT LONG VIDEO SUMMARY	1770
<i>Xiaopeng Ke, Boyu Chang, Hao Wu, Fengyuan Xu, Sheng Zhong, Nanjing University, China</i>	
IVMSP-6.3: CUT AND CONTINUOUS PASTE TOWARDS REAL-TIME DEEP FALL DETECTION	1775
<i>Sunhee Hwang, Minsong Ki, Seung-Hyun Lee, Sanghoon Park, Byoung-Ki Jeon, LG Uplus, Korea, Republic of</i>	
IVMSP-6.5: MANNET: A LARGE-SCALE MANIPULATED IMAGE DETECTION DATASET AND BASELINE EVALUATIONS	1780
<i>Aditya Singh, Saheb Chhabra, Puspita Majumdar, IIIT-Delhi, India; Richa Singh, Mayank Vatsa, IIT Jodhpur, India</i>	
IVMSP-6.6: APPROACHES TOWARD PHYSICAL AND GENERAL VIDEO ANOMALY DETECTION	1785
<i>Laura Kart, Niv Cohen, The Hebrew University of Jerusalem, Israel</i>	
IVMSP-7: PERCEPTION AND QUALITY MODELS	
IVMSP-7.1: CONSIDERING USER AGREEMENT IN LEARNING TO PREDICT THE AESTHETIC QUALITY	1790
<i>Suiyi Ling, Andreas Pastor, University of Nantes, France; Junle Wang, Tencent, China; Patrick Le Callet, Nantes Universite, Ecole Centrale Nantes, France</i>	
IVMSP-7.2: NO-REFERENCE QUALITY ASSESSMENT OF VARIABLE FRAME-RATE VIDEOS USING TEMPORAL BANDPASS STATISTICS	1795
<i>Qi Zheng, Yibo Fan, Xiaoyang Zeng, Fudan University, China; Zhengzhong Tu, Alan Bovik, University of Texas at Austin, United States of America</i>	
IVMSP-7.3: TOWARDS JOINT FRAME-LEVEL AND MOS QUALITY PREDICTIONS WITH LOW-COMPLEXITY OBJECTIVE MODELS	1800
<i>Joel Jung, Alexandre Giraud, Meijia Song, Songnan Li, Xiang Li, Shan Liu, Tencent Media Lab, United States of America</i>	
IVMSP-7.4: TEACHING CNNs TO MIMIC HUMAN VISUAL COGNITIVE PROCESS & REGULARISE TEXTURE-SHAPE BIAS	1805
<i>Satyam Mohla, Anshul Nasery, Biplab Banerjee, IIT Bombay, India</i>	

IVMSP-7.5: SUBJECTIVE AND OBJECTIVE QUALITY ASSESSMENT OF MOBILE GAMING VIDEO	1810
<i>Shaoguo Wen, Junle Wang, Ximing Chen, Yanqing Jing, Tencent, China; Suiyi Ling, University of Nantes, France; Patrick Le Callet, Nantes Université, Ecole Centrale Nantes, CNRS, LS2N, Nantes, France, France</i>	
IVMSP-7.6: ER-PIQA: A TASK-GUIDED PEDESTRIAN IMAGE QUALITY ASSESSMENT VIA EMBEDDING RECONSTRUCTION	1815
<i>Yanzhe Zhong, Bangjie Tang, Zhonggeng Liu, Yiming Zhu, Jun Yin, Advanced Research Institute of Zhejiang Dahua Technology Co.,Ltd., China; Huadong Pan, Advanced Research Institute of Zhejiang Dahua Technology Co.,Ltd., Zhejiang Provincial Key Laboratory of Harmonized Application of Vision & Transmission, China</i>	
IVMSP-8: APPLICATIONS I	
IVMSP-8.1: MULTISCALE CROWD COUNTING AND LOCALIZATION BY MULTITASK POINT SUPERVISION	1820
<i>Mohsen Zand, Haleh Damirchi, Andrew Farley, Mahdiyar Molahasani Majdabadi, Michael Greenspan, Ali Etemad, Queen's University, Canada</i>	
IVMSP-8.2: SUPER-RESOLUTION OF SATELLITE IMAGES BY TWO-DIMENSIONAL RRDB AND EDGE-ENHANCEMENT GENERATIVE ADVERSARIAL NETWORK	1825
<i>Yu-Zhang Chen, Tsung-Jung Liu, National Chung Hsing University, Taiwan; Kuan-Hsien Liu, National Taichung University of Science and Technology, Taiwan</i>	
IVMSP-8.3: LEVERAGING LOCAL TEMPORAL INFORMATION FOR MULTIMODAL SCENE CLASSIFICATION	1830
<i>Saurabh Sahu, Palash Goyal, Samsung Research America, United States of America</i>	
IVMSP-8.4: PREDICTING HUMAN MOTION USING KEY SUBSEQUENCES.....	1835
<i>Menghao Li, Mingtao Pei, Wei Liang, Beijing Institute of Technology, China</i>	
IVMSP-8.5: DYNAMIC TEXTURE RECOGNITION USING PDV HASHING AND DICTIONARY LEARNING ON MULTI-SCALE VOLUME LOCAL BINARY PATTERN	1840
<i>Ruxin Ding, Jianfeng Ren, Heng Yu, Jiawei Li, University of Nottingham Ningbo China, China</i>	
IVMSP-8.6: DO YOU LIVE A HEALTHY LIFE? ANALYZING LIFESTYLE BY VISUAL LIFE LOGGING	1845
<i>Qing Gao, Mingtao Pei, Hongyu Shen, Beijing Institute of technology, China</i>	
IVMSP-9: IMAGE AND VIDEO CODING II	
IVMSP-9.1: WEIGHTED WAVELET-BASED SPECTRAL-SPATIAL TRANSFORMS FOR CFA-SAMPLED RAW CAMERA IMAGE COMPRESSION CONSIDERING IMAGE FEATURES	1850
<i>Liping Huang, Taizo Suzuki, University of Tsukuba, Japan</i>	
IVMSP-9.2: JMPNET: JOINT MOTION PREDICTION FOR LEARNING-BASED VIDEO COMPRESSION	1855
<i>Dongyang Li, Zhenhong Sun, Zhiyu Tan, Xiuyu Sun, Fangyi Zhang, Yichen Qian, Hao Li, Alibaba Group, China, China</i>	
IVMSP-9.3: A LOW-PARAMETRIC MODEL FOR BIT-RATE ESTIMATION OF VVC RESIDUAL CODING	1860
<i>Fabian Brand, Christian Hergetz, André Kaup, Friedrich-Alexander-Universität Erlangen-Nürnberg, Germany</i>	
IVMSP-9.4: OPTE: ONLINE PER-TITLE ENCODING FOR LIVE VIDEO STREAMING.....	1865
<i>Vignesh V Menon, Hadi Amirpour, Christian Timmerer, University of Klagenfurt, Austria; Mohammad Ghanbari, University of Essex, United Kingdom of Great Britain and Northern Ireland</i>	

IVMSP-9.5: SADN: LEARNED LIGHT FIELD IMAGE COMPRESSION WITH SPATIAL-ANGULAR DECORRELATION	1870
<i>Kedeng Tong, Xin Jin, Chen Wang, Fan Jiang, Tsinghua University, China</i>	
IVMSP-9.6: HIERARCHICAL FEATURE AGGREGATION NETWORK FOR DEEP IMAGE COMPRESSION	1875
<i>Wenfeng Li, Zongcai Du, Hao He, Jie Tang, Gangshan Wu, Nanjing University, China</i>	
IVMSP-10: MACHINE LEARNING FOR IMAGE PROCESSING II	
IVMSP-10.1: ACCURATE INSTANCE SEGMENTATION VIA COLLABORATIVE LEARNING	1880
<i>Tianyou Chen, Xiaoguang Hu, Jin Xiao, Guofeng Zhang, Shaojie Wang, Beihang University, China</i>	
IVMSP-10.2: DYNAMIC BINARY NEURAL NETWORK BY LEARNING CHANNEL-WISE THRESHOLDS	1885
<i>Jiehua Zhang, Zhuo Su, Matti Pietikäinen, University of Oulu, China; Yanghe Feng, Xin Lu, National University of Defense Technology, China; Li Liu, 1. National University of Defense Technology; 2. University of Oulu, China</i>	
IVMSP-10.3: SELF-SUPERVISED LEARNING ON A LIGHTWEIGHT LOW-LIGHT IMAGE ENHANCEMENT MODEL WITH CURVE REFINEMENT	1890
<i>Wanyu Wu, Wei Wang, Xin Xu, Wuhan University of Science and Technology, China; Kui Jiang, Ruimin Hu, Wuhan University, China</i>	
IVMSP-10.4: SEMANTICALLY PROPORTIONAL PATCHMIX FOR FEW-SHOT LEARNING	1895
<i>Jingquan Wang, Jing Xu, Yu Pan, Harbin Institute of Technology (Shenzhen), China; Zenglin Xu, Harbin Institute of Technology (Shenzhen), PengCheng Laboratory, China</i>	
IVMSP-10.5: NOISE SUPPRESSION FOR IMPROVED FEW-SHOT LEARNING	1900
<i>Zhikui Chen, Tiandong Ji, Suhua Zhang, Fangming Zhong, Dalian University of Technology, China</i>	
IVMSP-10.6: ONLINE CONTINUAL LEARNING USING ENHANCED RANDOM VECTOR FUNCTIONAL LINK NETWORKS	1905
<i>Cheryl Sze Yin Wong, Yang Guo, ArulMurugan Ambikapathi, Savitha Ramasamy, Institute for Infocomm Research, Singapore</i>	
IVMSP-11: VIDEO PROCESSING	
IVMSP-11.1: A GENERALIZED KERNEL RISK SENSITIVE LOSS FOR ROBUST TWO-DIMENSIONAL SINGULAR VALUE DECOMPOSITION	1910
<i>Miaohua Zhang, Yongsheng Gao, Jun Zhou, Griffith University, Australia</i>	
IVMSP-11.2: VIDEO FRAME INTERPOLATION VIA LOCAL LIGHTWEIGHT BIDIRECTIONAL ENCODING WITH CHANNEL ATTENTION CASCADE	1915
<i>Xiangling Ding, Hunan University of Science and Technology, China; Pu Huang, Dengyong Zhang, Changsha University of Science and Technology, China; Xianfeng Zhao, Chinese Academy of Sciences, China</i>	
IVMSP-11.3: SAIN: SIMILARITY-AWARE VIDEO FRAME INTERPOLATION	1920
<i>Yue Lv, Wenming Yang, Qingmin Liao, Tsinghua University, China; Wangmeng Zuo, Harbin Institute of Technology, China; Rui Zhu, City, University of London, United Kingdom of Great Britain and Northern Ireland</i>	
IVMSP-11.4: SELF-LEARNED VIDEO SUPER-RESOLUTION WITH AUGMENTED SPATIAL AND TEMPORAL CONTEXT	1925
<i>Zeja Fan, Jiaying Liu, Wenhan Yang, Zongming Guo, Wangxuan Institute of Computer Technology, Peking University, China; Wei Xiang, Bigo Technology PTE. LTD, China</i>	

IVMSP-11.5: DEFORMABLE CONVOLUTION DENSE NETWORK FOR COMPRESSED VIDEO QUALITY ENHANCEMENT	1930
<i>Jiahui Liu, Mingcai Zhou, Meng Xiao, Alibaba Group, China</i>	
IVMSP-11.6: CONVOLUTIONAL ISTA NETWORK WITH TEMPORAL CONSISTENCY CONSTRAINTS FOR VIDEO RECONSTRUCTION FROM EVENT CAMERAS	1935
<i>Siyang Liu, Roxana Alexandru, Pier Luigi Dragotti, Imperial College London, United Kingdom of Great Britain and Northern Ireland</i>	
IVMSP-12: IMAGE AND VIDEO UNDERSTANDING II	
IVMSP-12.1: PMP-NET: RETHINKING VISUAL CONTEXT FOR SCENE GRAPH GENERATION	1940
<i>Xuezhi Tong, College of Intelligence and Computing, Tianjin University, China; Rui Wang, Zhejiang Lab and State Key Laboratory of Information Security, Institute of Information Engineering, Chinese Academy of Sciences, China; Chuan Wang, Sanyi Zhang, Xiaochun Cao, State Key Laboratory of Information Security, Institute of Information Engineering, Chinese Academy of Sciences, China</i>	
IVMSP-12.2: IMPROVE IMAGE CAPTIONING VIA RELATION MODELING	1945
<i>Feicheng Huang, Zhixin Li, Guangxi Normal University, China</i>	
IVMSP-12.3: EQUAL LOSS: A SIMPLE LOSS FUNCTION FOR NOISE ROBUST LEARNING	1950
<i>Lei Cui, Tsinghua University, China; Huan Peng, Huazhong University of Science and Technology, China; Yangguang Li, Chuming Li, Sensetime Group, China; Xingrun Xing, Beihang University, China</i>	
IVMSP-12.4: INFORMATIVE ATTENTION SUPERVISION FOR GROUNDED VIDEO DESCRIPTION	1955
<i>Boyang Wan, Wenhui Jiang, Yuming Fang, Jiangxi University of Finance and Economics, China</i>	
IVMSP-12.5: SPATIAL-CONTEXT-AWARE DEEP NEURAL NETWORK FOR MULTI-CLASS IMAGE CLASSIFICATION	1960
<i>Jialu Zhang, Qian Zhang, Jianfeng Ren, University of Nottingham Ningbo China, China; Yitian Zhao, Chinese Academy of Sciences, China; Jiang Liu, Southern University of Science and Technology, China</i>	
IVMSP-12.6: TRANSTL: SPATIAL-TEMPORAL LOCALIZATION TRANSFORMER FOR MULTI-LABEL VIDEO CLASSIFICATION	1965
<i>Hongjun Wu, Mengzhu Li, Hongzhe Liu, Cheng Xu, Xuwei Li, Beijing Union University, China; Yongcheng Liu, Chinese Academy of Sciences, China</i>	
IVMSP-13: MULTI-MODAL ANALYSIS AND SYNTHESIS	
IVMSP-13.1: DEEP VIDEO INPAINTING GUIDED BY AUDIO-VISUAL SELF-SUPERVISION	1970
<i>Kyuyeon Kim, Junsik Jung, Woo Jae Kim, Sung-Eui Yoon, Korea Advanced Institute of Science and Technology, Korea, Republic of</i>	
IVMSP-13.2: NAVIGATING AUDIO-VISUAL EVENT DETECTION ACROSS MISMATCHED MODALITIES	1975
<i>Guangwei Li, Xuenan Xu, Mengyue Wu, Kai Yu, Shanghai Jiao Tong University, China</i>	
IVMSP-13.3: LOOK, LISTEN AND PAY MORE ATTENTION: FUSING MULTI-MODAL INFORMATION FOR VIDEO VIOLENCE DETECTION	1980
<i>Dong-Lai Wei, Yang Liu, Jing Liu, Xin-Hua Zeng, Fudan University, China; Chen-Geng Liu, The University of Melbourne, China; Xiao-Guang Zhu, Shanghai Jiao Tong University, China</i>	

IVMSP-13.4: MULTI-MODAL LEARNING WITH TEXT MERGING FOR TEXTVQA 1985
Changsheng Xu, Zhenlong Xu, Yifan He, Shuigeng Zhou, Fudan Univeristy, China; Jihong Guan, Tongji University, China

IVMSP-13.6: A NOVEL PART FEATURE INTEGRATION AND FUSION METHOD FOR 1990
FINE-GRAINED VEHICLE RECOGNITION
Ping Wang, Yijie Cao, Lei Lu, Xi'an Jiaotong University, China

IVMSP-14: OBJECT DETECTION II

IVMSP-14.1: MONOCULAR VEHICLE 3D BOUNDING BOX ESTIMATION USING 1995
HOMOGRAHY AND GEOMETRY IN TRAFFIC SCENE
Yiqiang Chen, Feng Liu, Ke Pei, Huawei Technology, China

IVMSP-14.2: FSM: FEATURE SAMPLING MODULE FOR OBJECT DETECTION 2000
Xin Yi, Bo Ma, JiaHao Wu, Beijing Institute of Technology, China

IVMSP-14.3: RETHINKING TWO-B-REAL NET FOR REAL-TIME SALIENT OBJECT 2005
DETECTION
Senyun Kuang, Shijin Meng, Bo Xiao, Southwest Jiaotong University, China; Lv Tang, Nanjing University, China; Bo Li, Youtu Lab, Tencent, China; , ,

IVMSP-14.4: BALANCED RANKING AND SORTING FOR CLASS INCREMENTAL 2010
OBJECT DETECTION
Bo Cui, Shan Yu, CASIA, China; Hui Qu, Xuhui Huang, X-Lab, the Second Academy of CASIC, China

IVMSP-14.5: MULTI-SCALE REINFORCEMENT LEARNING STRATEGY FOR OBJECT 2015
DETECTION
Yihao Luo, Xiang Cao, Juntao Zhang, Leixilan Pan, Tianjiang Wang, Qi Feng, Huazhong University of Science and Technology, China

IVMSP-14.6: DEEP OBJECT DETECTION WITH EXAMPLE ATTRIBUTE BASED 2020
PREDICTION MODULATION
Zhihao Wu, Chengliang Liu, Chao Huang, Jie Wen, Yong Xu, Harbin Institute of Technology, Shenzhen, China

IVMSP-15: IMAGE COMMUNICATION AND QUALITY MODELS

IVMSP-15.1: UNIVERSAL EFFICIENT VARIABLE-RATE NEURAL IMAGE 2025
COMPRESSION
Shanzhi Yin, Chao Li, Youneng Bao, Yongsheng Liang, Harbin Institute of Technology, Shenzhen, China; Fanyang Meng, Wei Liu, Peng Cheng Laboratory, China

IVMSP-15.2: ADDERIC: TOWARDS LOW COMPUTATION COST IMAGE 2030
COMPRESSION
Bowen Li, Chao Li, Youneng Bao, Yongsheng Liang, Harbin Institute of Technology, Shenzhen, China, China; Yao Xin, Fanyang Meng, Peng Cheng Laboratory, Shenzhen, China, China

IVMSP-15.3: DCNGAN: A DEFORMABLE CONVOLUTION-BASED GAN WITH QP 2035
ADAPTATION FOR PERCEPTUAL QUALITY ENHANCEMENT OF COMPRESSED VIDEO
Saiping Zhang, Fuzheng Yang, Xidian University, China; Luis Herranz, Universitat Autònoma de Barcelona, Spain; Marta Mrak, Marc Gorriz Blanch, British Broadcasting Corporation, United Kingdom of Great Britain and Northern Ireland; Shuai Wan, Northwestern Polytechnical University, China

IVMSP-15.4: SPECIALISED VIDEO QUALITY MODEL FOR ENHANCED USER 2040
GENERATED CONTENT (UGC) WITH SPECIAL EFFECTS
Anne-Flore Perrin, Yiejing Xie, Nantes Université, Ecole Centrale Nantes, CAPACITES SAS, CNRS, LS2N, UMR 6004, F-44000 Nantes, France, France; Tao Zhang, Yiting Liao, Junlin Li, ByteDance Inc., United States of America; Patrick Le Callet, Nantes Université, Ecole Centrale Nantes, CNRS, LS2N, UMR 6004, F-44000 Nantes, France

IVMSP-15.5: IMPROVING MAXIMUM LIKELIHOOD DIFFERENCE SCALING METHOD TO MEASURE INTER CONTENT SCALE	2045
<i>Andréas Pastor, Patrick Le Callet, University of Nantes, France; Lukáš Krasula, Xiaoqing Zhu, Zhi Li, Netflix Inc., United States of America</i>	
IVMSP-15.6: TEXTURE INFORMATION BOOSTS VIDEO QUALITY ASSESSMENT	2050
<i>Ao-Xiang Zhang, Yuan-Gen Wang, Guangzhou University, China</i>	
IVMSP-16: APPLICATIONS II	
IVMSP-16.1: PLUG-AND-PLAY AND RELAY REGULARIZATIONS ON NOISY LOW RANK TENSOR COMPLETION FOR SNAPSHOT MULTISPECTRAL IMAGE RESTORATION	2055
<i>Keisuke Ozawa, DENSO IT Laboratory, Japan</i>	
IVMSP-16.2: LERPS: LIGHTING ESTIMATION AND RELIGHTING FOR PHOTOMETRIC STEREO	2060
<i>Ashish Tiwari, Shanmuganathan Raman, IIT Gandhinagar, India</i>	
IVMSP-16.3: A UNIFIED TWO-STAGE MODEL FOR SEPARATING SUPERIMPOSED IMAGES	2065
<i>Huiyu Duan, Xiongkuo Min, Wei Shen, Guangtao Zhai, Shanghai Jiao Tong University, China</i>	
IVMSP-16.4: PARAMETER-FREE STYLE PROJECTION FOR ARBITRARY IMAGE STYLE TRANSFER	2070
<i>Siyu Huang, Harvard University, China; Haoyi Xiong, Qingzhong Wang, Zeyu Chen, Dejing Dou, Baidu Research, China; Tianyang Wang, Austin Peay State University, United States of America; Bihan Wen, Nanyang Technological University, Singapore; Jun Huan, Amazon, China</i>	
IVMSP-16.5: OPTIMIZATION OF COMPRESSIVE LIGHT FIELD DISPLAY IN DUAL-GUIDED LEARNING	2075
<i>Yangfan Sun, Zhu Li, UMKC, United States of America; Li Li, University of Science and Technology of China, United States of America; Shizheng Wang, Chinese Academy of Sciences, United States of America; Wei Gao, Peking University, United States of America</i>	
IVMSP-16.6: ARM 4-BIT PQ: SIMD-BASED ACCELERATION FOR APPROXIMATE NEAREST NEIGHBOR SEARCH ON ARM	2080
<i>Yusuke Matsui, The University of Tokyo, Japan; Yoshiki Imaizumi, Naoya Miyamoto, Naoki Yoshifuji, Fixstars Corporation, Japan</i>	
IVMSP-17: ENHANCEMENT AND RESTORATION I	
IVMSP-17.1: ITERATIVE LEARNING FOR DISTORTED IMAGE RESTORATION	2085
<i>Chao Wang, Yi Gu, Jie Li, Xinlei He, Zirui Zhang, Chentao Wu, Shanghai Jiao Tong University, China; Yuting Gao, Texas A&M University, United States of America</i>	
IVMSP-17.2: JE2NET: JOINT EXPLOITATION AND EXPLORATION IN REINFORCEMENT LEARNING BASED IMAGE RESTORATION	2090
<i>Xiaoyu Zhang, Wei Gao, Ge Li, School of Electronic and Computer Engineering, Peking University, Shenzhen, China. Peng Cheng Laboratory, Shenzhen, China., China; Hui Yuan, School of Control Science and Engineering, Shandong University, Jinan, China, China</i>	
IVMSP-17.3: MULTIPLE PATCH-AWARE NETWORK FOR FASTER REAL-WORLD IMAGE DEHAZING	2095
<i>Kun Yang, Juan Zhang, Xiaoqi Lang, Shanghai University of Engineering Science, China</i>	
IVMSP-17.4: LEARNING TO FUSE HETEROGENEOUS FEATURES FOR LOW-LIGHT IMAGE ENHANCEMENT	2100
<i>Zhenyu Tang, Long Ma, Xiaoke Shang, Xin Fan, Dalian University of Technology, China</i>	

IVMSP-17.5: DEEP SCALE-AWARE IMAGE SMOOTHING	2105
<i>Jiachun Li, Kunkun Qin, Ruotao Xu, South China University of Technology, China; Hui Ji, National University of Singapore, Singapore</i>	
IVMSP-17.6: A MULTISCALE GRADIENT-BACKPROPAGATION OPTIMIZATION	2110
FRAMEWORK FOR DEFORMABLE CONVOLUTION BASED COMPRESSED VIDEO ENHANCEMENT	
<i>Yanbo Gao, Shuai Li, Xun Cai, Shandong University, China; Menghu Jia, Mao Ye, University of Electronic Science and Technology of China, China; Frédéric Dufaux, Université Paris-Saclay, France</i>	
IVMSP-18: MACHINE LEARNING FOR IMAGE PROCESSING III	
IVMSP-18.1: DOWNSTREAM AUGMENTATION GENERATION FOR CONTRASTIVE	2115
LEARNING	
<i>Tomohiro Hayase, Suguru Yasutomi, Fujitsu, Ltd., Japan; Nakamasa Inoue, Tokyo Institute of Technology, Japan</i>	
IVMSP-18.2: FEW-SHOT LEARNING WITH IMPROVED LOCAL REPRESENTATIONS	2120
VIA BIAS RECTIFY MODULE	
<i>Chao Dong, Qi Ye, Wenchao Meng, Kaixiang Yang, Zhejiang University, China</i>	
IVMSP-18.3: IMAGE-TO-VIDEO RE-IDENTIFICATION VIA MUTUAL DISCRIMINATIVE	2125
KNOWLEDGE TRANSFER	
<i>Pichao Wang, Fan Wang, Hao Li, Alibaba Group, United States of America</i>	
IVMSP-18.4: DYNSSNN: A DYNAMIC APPROACH TO REDUCE REDUNDANCY IN	2130
SPIKING NEURAL NETWORKS	
<i>Fangxin Liu, Wenbo Zhao, Yongbiao Chen, Zongwu Wang, Shanghai Jiaotong University, China; Dai Fei, Academy of System Engineering, Academy of Military Sciences, China</i>	
IVMSP-18.5: MEJIGCLU: MORE EFFECTIVE JIGSAW CLUSTERING FOR	2135
UNSUPERVISED VISUAL REPRESENTATION LEARNING	
<i>Yongsheng Zhang, Qing Liu, Yang Zhao, Yixiong Liang, Central South University, China</i>	
IVMSP-18.6: GANET: UNARY ATTENTION REACHES PAIRWISE ATTENTION VIA	2140
IMPLICIT GROUP CLUSTERING IN LIGHT-WEIGHT CNNs	
<i>Cheng Zhuang, Yunlian Sun, Nanjing University of Science and Technology, China</i>	
IVMSP-19: SUPER-RESOLUTION	
IVMSP-19.1: FIND THE WAY BACK: INVERTIBLE KERNEL ESTIMATOR FOR	2145
BLIND IMAGE SUPER-RESOLUTION	
<i>Ting-Wei Chang, Wei-Chen Chiu, Ching-Chun Huang, National Yang Ming Chiao Tung University, Taiwan, Province of China</i>	
IVMSP-19.3: FINE-GRAINED DYNAMIC LOSS FOR ACCURATE SINGLE-IMAGE	2150
SUPER-RESOLUTION	
<i>Haoquan Wang, Gang Zhang, Tianjin university, China; Zhichun Lei, University of Applied Science Ruhr West (HRW) in Germany, Germany</i>	
IVMSP-19.4: MULTI-FRAME SUPER-RESOLUTION WITH RAW IMAGES VIA	2155
MODIFIED DEFORMABLE CONVOLUTION	
<i>Gongzhe Li, Linwei Qiu, Haopeng Zhang, Fengying Xie, Zhiguo Jiang, Beihang University, China</i>	
IVMSP-19.5: LOCAL-GLOBAL FEATURE AGGREGATION FOR LIGHT FIELD IMAGE	2160
SUPER-RESOLUTION	
<i>Yan Wang, Yao Lu, Shunzhou Wang, Wenyaoy Zhang, Zijian Wang, Beijing Institute of Technology, China</i>	

IVMSP-19.6: PYRAMID FUSION ATTENTION NETWORK FOR SINGLE IMAGE SUPER-RESOLUTION	2165
<i>Hao He, Zongcai Du, Wenfeng Li, Jie Tang, Gangshan Wu, Nanjing University, China</i>	
IVMSP-20: ACTION RECOGNITION	
IVMSP-20.1: VCD: VIEW-CONSTRAINT DISENTANGLEMENT FOR ACTION RECOGNITION	2170
<i>Xian Zhong, Zhuo Zhou, Wenxuan Liu, Xuemei Jia, Wuhan University of Technology, China; Kui Jiang, Zheng Wang, Wuhan University, China; Wenxin Huang, Hubei University, China</i>	
IVMSP-20.2: PRIVACY-PRESERVING ACTION RECOGNITION	2175
<i>Chengming Zou, Hubei Key Laboratory of Transportation Internet of Things, Wuhan University of Technology, Wuhan, 430070, China, China; Ducheng Yuan, School of Computer Science and Technology, Wuhan University of Technology, Wuhan, 430070, China, China; Long Lan, Haoang Chi, Institute for Quantum Information & State Key Laboratory of High Performance Computing, College of Computer, National University of Defense Technology, Changsha 410073, China, China</i>	
IVMSP-20.3: SPATIO-TEMPORAL MOTION AGGREGATION NETWORK FOR VIDEO ACTION DETECTION	2180
<i>Hongcheng Zhang, Xu Zhao, Shanghai Jiao Tong University, China</i>	
IVMSP-20.4: TP-VIT: A TWO-PATHWAY VISION TRANSFORMER FOR VIDEO ACTION RECOGNITION	2185
<i>Yanhao Jing, Feng Wang, East China Normal University, China</i>	
IVMSP-20.5: LEARNING TASK-SPECIFIC REPRESENTATION FOR VIDEO ANOMALY DETECTION WITH SPATIAL-TEMPORAL ATTENTION	2190
<i>Yang Liu, Jing Liu, Donglai Wei, Xiaohong Huang, Liang Song, Fudan University, China; Xiaoguang Zhu, Shanghai Jiao Tong University, China</i>	
IVMSP-20.6: W-ART: ACTION RELATION TRANSFORMER FOR WEAKLY-SUPERVISED TEMPORAL ACTION LOCALIZATION	2195
<i>Mengzhu Li, Hongjun Wu, Hongzhe Liu, Cheng Xu, Xuwei Li, Beijing Union University, China; Yongcheng Liu, Chinese Academy of Sciences, China</i>	
IVMSP-21: ATTENTION FOR IMAGE AND VIDEO PROCESSING	
IVMSP-21.1: MS-ROCANET: MULTI-SCALE RESIDUAL ORTHOGONAL-CHANNEL ATTENTION NETWORK FOR SCENE TEXT DETECTION	2200
<i>Jinpeng Liu, Song Wu, DeHong He, Guoqiang Xiao, SouthWest University, China</i>	
IVMSP-21.2: BI-DIRECTIONAL NORMALIZATION AND COLOR ATTENTION-GUIDED GENERATIVE ADVERSARIAL NETWORK FOR IMAGE ENHANCEMENT	2205
<i>Shan Liu, Guoqiang Xiao, Xiaohui Xu, Song Wu, Southwest University, China</i>	
IVMSP-21.3: DUAL-ATTENTION NETWORK FOR FEW-SHOT SEGMENTATION	2210
<i>Zhikui Chen, Han Wang, Suhua Zhang, Fangming Zhong, Dalian University of Technology, China</i>	
IVMSP-21.4: ATTENTION GUIDED INVARIANCE SELECTION FOR LOCAL FEATURE DESCRIPTORS	2215
<i>Jiapeng Li, Ge Li, Peking University, China; Thomas H Li, Advanced Institute of Information Technology, Peking University, China</i>	
IVMSP-21.5: ATTENTION PROBE: VISION TRANSFORMER DISTILLATION IN THE WILD	2220
<i>Jiahao Wang, Mingdeng Cao, Shuwei Shi, Yujiu Yang, Tsinghua Shenzhen International Graduate School, China; Baoyuan Wu, The Chinese University of Hong Kong, Shenzhen, China</i>	

IVMSP-21.6: STACKED MULTI-SCALE ATTENTION NETWORK FOR IMAGE COLORIZATION	2225
<i>Bin Jiang, Fangqiang Xu, Jun Xia, Chao Yang, Wei Huang, Yun Huang, Hunan university, China</i>	
IVMSP-22: OBJECT DETECTION I	
IVMSP-22.1: CRPN: DISTINGUISH NOVEL CATEGORIES VIA CLASS-RELEVANT REGION PROPOSAL NETWORK FOR FEW-SHOT OBJECT DETECTION	2230
<i>Han Wang, Yali Li, Shengjin Wang, Tsinghua University, China</i>	
IVMSP-22.2: AN EFFICIENT FRAMEWORK FOR DETECTION AND RECOGNITION OF NUMERICAL TRAFFIC SIGNS	2235
<i>Zhishan Li, Lei Xie, Hongye Su, Zhejiang University, China; Mingmu Chen, Yifan He, Reconova Technologies Co., Ltd, China</i>	
IVMSP-22.3: DIVERGENCE-GUIDED FEATURE ALIGNMENT FOR CROSS-DOMAIN OBJECT DETECTION	2240
<i>Zongyao Li, Ren Togo, Takahiro Ogawa, Miki Haseyama, Hokkaido University, Japan</i>	
IVMSP-22.4: PGTRNET: TWO-PHASE WEAKLY SUPERVISED OBJECT DETECTION WITH PSEUDO GROUND TRUTH REFINEMENT	2245
<i>Jun Wang, University of Warwick, United Kingdom of Great Britain and Northern Ireland; Hefeng Zhou, Shanghai Jiao Tong University, China; Xiaohan Yu, School of Engineering and Built Environment, Australia</i>	
IVMSP-22.5: NOVEL INSTANCE MINING WITH PSEUDO-MARGIN EVALUATION FOR FEW-SHOT OBJECT DETECTION	2250
<i>Weijie Liu, Chong Wang, Shenghao Yu, Chenchen Tao, Ningbo University, China; Jun Wang, China University of Mining and Technology, China; Jiafei Wu, SenseTime Research, China</i>	
IVMSP-22.6: BIP-NET: BIDIRECTIONAL PERSPECTIVE STRATEGY BASED ARBITRARY-SHAPED TEXT DETECTION NETWORK	2255
<i>Chuang Yang, Mulin Chen, Yuan Yuan, Qi Wang, School of Computer Science, and School of Artificial Intelligence, Optics and Electronics (iOPEN), Northwestern Polytechnical University, Xi'an 710072, Shaanxi, P. R. China, China</i>	
IVMSP-23: DEPTH ESTIMATION & POINT CLOUD PROCESSING	
IVMSP-23.1: A NOVEL LIGHTWEIGHT NETWORK FOR FAST MONOCULAR DEPTH ESTIMATION	2260
<i>Tim Heydrich, Yimin Yang, Lakehead University, Canada; Xiangyu Ma, Yu Liu, Tianjin University, China; Shan Du, The University of British Columbia, Canada</i>	
IVMSP-23.2: A LIGHTWEIGHT SELF-SUPERVISED TRAINING FRAMEWORK FOR MONOCULAR DEPTH ESTIMATION	2265
<i>Tim Heydrich, Yimin Yang, Lakehead University, Canada; Shan Du, The University of British Columbia, Canada</i>	
IVMSP-23.3: PU-REFINER: A GEOMETRY REFINER WITH ADVERSARIAL LEARNING FOR POINT CLOUD UPSAMPLING	2270
<i>Hao Liu, Hui Yuan, Shuai Li, Shandong University, China; Hamzaoui Raouf, De Montfort University, United Kingdom of Great Britain and Northern Ireland; Wei Gao, Peking University, China</i>	
IVMSP-23.4: CF-NET: COMPLEMENTARY FUSION NETWORK FOR ROTATION INVARIANT POINT CLOUD COMPLETION	2275
<i>Bo-Fan Chen, Yang-Ming Yeh, Yi-Chang Lu, National Taiwan University, Taiwan</i>	
IVMSP-23.5: TH-NET: A METHOD OF SINGLE 3D OBJECT TRACKING BASED ON TRANSFORMERS AND HAUSDORFF DISTANCE	2280
<i>Zihao Zhang, Nan Sang, Xupeng Wang, University of Electronic Science and Technology of China, China</i>	

IVMSP-23.6: ENRICH FEATURES FOR FEW-SHOT POINT CLOUD CLASSIFICATION	2285
<i>Hengxin Feng, Weifeng Liu, Yanjiang Wang, Baodi Liu, China University of Petroleum (East China), China</i>	
IVMSP-24: APPLICATIONS III	
IVMSP-24.1: SEMI-SUPERVISED 360° DEPTH ESTIMATION FROM MULTIPLE FISHEYE CAMERAS WITH PIXEL-LEVEL SELECTIVE LOSS	2290
<i>Jaewoo Lee, Daeul Park, Dongwook Lee, Daehyun Ji, Samsung Advanced Institute of Technology, Korea, Republic of</i>	
IVMSP-24.2: UNDERWATER STEREO MATCHING VIA UNSUPERVISED APPEARANCE AND FEATURE ADAPTATION NETWORKS	2295
<i>Wei Zhong, Yazhi Yuan, Xinchun Ye, Dian Zheng, Rui Xu, Dalian University of Technology, China</i>	
IVMSP-24.3: DOMAIN ADAPTATION VIA MUTUAL INFORMATION MAXIMIZATION FOR HANDWRITING RECOGNITION	2300
<i>Pei Tang, Liangrui Peng, Ruijie Yan, Haodong Shi, Gang Yao, Changsong Liu, Tsinghua University, China; Jie Li, Yuqi Zhang, Shanghai Pudong Development Bank, China</i>	
IVMSP-24.4: ATTRIBUTE-CONDITIONED FACE SWAPPING NETWORK FOR LOW-RESOLUTION IMAGES	2305
<i>Ang Li, Chilin Fu, Xiaolu Zhang, Jun Zhou, Ant Group, China; Jian Hu, Queen mary university of London, China</i>	
IVMSP-24.5: LEARNING MULTIPLE EXPLAINABLE AND GENERALIZABLE CUES FOR FACE ANTI-SPOOFING	2310
<i>Ying Bian, Peng Zhang, Jingjing Wang, Chunmao Wang, Shiliang Pu, Hikvision Research Institute, China, China</i>	
IVMSP-24.6: OFF-THE-GRID COVARIANCE-BASED SUPER-RESOLUTION FLUCTUATION MICROSCOPY	2315
<i>Bastien Laville, Laure Blanc-Féraud, Gilles Aubert, Université Côte d'Azur, France</i>	
IVMSP-25: ENHANCEMENT AND RESTORATION II	
IVMSP-25.1: SIMULTANEOUS NONLOCAL LOW-RANK AND DEEP PRIORS FOR POISSON DENOISING	2320
<i>Zhiyuan Zha, Bihan Wen, Nanyang Technological University, Singapore; Xin Yuan, Nokia Bell Labs, United States of America; Jiantao Zhou, University of Macau, China; Ce Zhu, University of Electronic Science and Technology of China, China</i>	
IVMSP-25.2: DOUBLE CLOSED-LOOP NETWORK FOR IMAGE DEBLURRING.....	2325
<i>Yiming Liu, Hunan Normal University, China; Yanni Zhang, Qiang Li, Jun Kong, Jianzhong Wang, Northeast Normal University, China; Miao Qi, Changchun Humanities and Sciences College, China</i>	
IVMSP-25.3: SINGLE IMAGE DE-RAINING WITH HIGH-LOW FREQUENCY GUIDANCE	2330
<i>Ying Zhang, Youjun Xiang, Lei Cai, Yuli Fu, Wanliang Huo, Junjun Xia, South China University of Technology, China</i>	
IVMSP-25.4: DETAIL GENERATION AND FUSION NETWORKS FOR IMAGE INPAINTING	2335
<i>Wu Yang, Wuzhen Shi, Shenzhen University, China</i>	
IVMSP-25.6: ADAPTIVE WEIGHTED NETWORK WITH EDGE ENHANCEMENT MODULE FOR MONOCULAR SELF-SUPERVISED DEPTH ESTIMATION	2340
<i>Hong Liu, Ying Zhu, Guoliang Hua, Weiho Huang, Runwei Ding, Key Laboratory of Machine Perception, Shenzhen Graduate School, Peking University, China</i>	
IVMSP-25.7: PAS-MEF: MULTI-EXPOSURE IMAGE FUSION BASED ON PRINCIPAL COMPONENT ANALYSIS, ADAPTIVE WELL-EXPOSEDNESS AND SALIENCY MAP	2345
<i>Diclehan Karakaya, Oguzhan Ulucan, Mehmet Turkan, Izmir University of Economics, Turkey</i>	

IVMSP-26: MACHINE LEARNING FOR IMAGE PROCESSING IV

IVMSP-26.1: PDD-NET: A PRECISE DEFECT DETECTION NETWORK BASED ON 2350 POINT SET REPRESENTATION

Miaoju Ban, Runwei Ding, Jian Zhang, Tianyu Guo, Tao Wang, Key Laboratory of Machine Perception, Peking University, Shenzhen Graduate School, China

IVMSP-26.2: SOLVING THE LONG-TAILED PROBLEM VIA INTRA- AND 2355 INTER-CATEGORY BALANCE

Renhui Zhang, Tiancheng Lin, Rui Zhang, Yi Xu, Shanghai Jiao Tong University, China

IVMSP-26.3: EXTRACTING AND DISTILLING DIRECTION-ADAPTIVE KNOWLEDGE 2360 FOR LIGHTWEIGHT OBJECT DETECTION IN REMOTE SENSING IMAGES

Zhanchao Huang, Wei Li, Ran Tao, Beijing Institute of Technology, China

IVMSP-26.4: PSEUDO-INTERACTING GUIDED NETWORK FOR FEW-SHOT 2365 SEGMENTATION

Xiaoliu Luo, Zhao Duan, Jin Tan, Taiping Zhang, Chongqing University, China; Jing Luo, Zhengzhou University, China

IVMSP-26.5: FEW-SHOT GENERATION BY MODELING STEREOSCOPIC PRIORS 2370

Yuehui Wang, Qing Wang, Dongyu Zhang, Sun-yat sen university, China

IVMSP-26.6: RELATIVE VIEWPOINT ESTIMATION BASED ON STRUCTURED 3D 2375 REPRESENTATION ALIGNMENT

Kohei Matsuzaki, Kei Kawamura, KDDI Research, Inc., Japan

IVMSP-27: SEGMENTATION

IVMSP-27.1: DEEP MARKOV CLUSTERING FOR PANOPTIC SEGMENTATION 2380

Minxiang Ye, YiFei Zhang, Shiqiang Zhu, Anhuan Xie, Dan Zhang, Zhejiang Lab, China

IVMSP-27.2: MULTI-TASK LEARNING IMPROVES THE BRAIN STROKE LESION 2385 SEGMENTATION

Libo Liu, Qingmao Hu, Shenzhen Institutes of Advanced Technology, Chinese Academy of Sciences, Shenzhen, China; School of Artificial Intelligence, University of Chinese Academy of Sciences, Beijing, China, China; Chengjian Huang, Shenzhen Institutes of Advanced Technology, Chinese Academy of Sciences, Shenzhen, China; University of Chinese Academy of Sciences, Beijing, China, China; Chunsheng Cai, Department of Cerebrovascular Disease, Huizhou First People's Hospital, Huizhou, China, China; Xiaodong Zhang, Shenzhen Institutes of Advanced Technology, Chinese Academy of Sciences, Shenzhen, China, China

IVMSP-27.3: MIXED TRANSFORMER U-NET FOR MEDICAL IMAGE SEGMENTATION 2390

Hongyi Wang, Shiao Xie, Lanfen Lin, Ruofeng Tong, Zhejiang University, China; Yutaro Iwamoto, Yen-Wei Chen, Ritsumeikan University, Japan; Xian-Hua Han, Yamaguchi University, Japan

IVMSP-27.4: CONTRASTIVE TRANSLATION LEARNING FOR MEDICAL IMAGE 2395 SEGMENTATION

Wankang Zeng, Wenkang Fan, Dongfang Shen, Yinran Chen, Xiongbiao Luo, Xiamen University, China

IVMSP-27.5: FAST VIDEO OBJECT SEGMENTATION VIA DYNAMIC YOLACT 2400

Tianfang Meng, Wenqiang Zhang, Fudan University, China

IVMSP-27.6: DEPTH REMOVAL DISTILLATION FOR RGB-D SEMANTIC 2405 SEGMENTATION

Tiyu Fang, Zhen Liang, Zihao Dong, Jinping Li, University of Jinan, China; Xiuli Shao, Nankai University, China

IVMSP-28: FACES

IVMSP-28.1: MASK-BASED ATTENTION PARALLEL NETWORK FOR IN-THE-WILD FACIAL EXPRESSION RECOGNITION 2410

Lingzhao Ju, Xu Zhao, Shanghai Jiao Tong University, Department of Automation, Shanghai, China, China

IVMSP-28.2: SDNET: LIGHTWEIGHT FACIAL EXPRESSION RECOGNITION FOR SAMPLE DISEQUILIBRIUM 2415

Lifang Zhou, Siqin Li, Yi Wang, Junlin Liu, Chongqing University of Posts and Telecommunications, China

IVMSP-28.3: A NOVEL MICRO-EXPRESSION RECOGNITION APPROACH USING ATTENTION-BASED MAGNIFICATION-ADAPTIVE NETWORKS 2420

Wei Mengting, Zheng Wenming, Zong Yuan, Jiang Xingxun, Lu Cheng, Liu Jiateng, Southeast University, China

IVMSP-28.4: LIPREADING MODEL BASED ON WHOLE-PART COLLABORATIVE LEARNING 2425

Weidong Tian, Housen Zhang, Chen Peng, Zhong-Qiu Zhao, Hefei University of Technology, China

IVMSP-28.5: WHAT IS THE PATIENT LOOKING AT? ROBUST GAZE-SCENE INTERSECTION UNDER FREE-VIEWING CONDITIONS 2430

Ahmed Al-Hindawi, Marcela Vizcaychipi, Yiannis Demiris, Imperial College London, United Kingdom of Great Britain and Northern Ireland

IVMSP-28.6: GAZEATTENTIONNET: GAZE ESTIMATION WITH ATTENTIONS..... 2435

Haoxian Huang, Luqian Ren, Zhuo Yang, Yinwei Zhan, Guangdong University of Technology, China; Qieshi Zhang, Chinese Academy of Sciences, China; Jujian Lv, Guangdong Polytechnic Normal University, China

IVMSP-29: INVERSE PROBLEMS

IVMSP-29.1: LOW-LIGHT IMAGE ENHANCEMENT VIA FEATURE RESTORATION 2440

Yang Yang, Xiaojie Guo, Tianjin university, China; Yonghua Zhang, Henan University, China

IVMSP-29.2: HIRL: HYBRID IMAGE RESTORATION BASED ON HIERARCHICAL DEEP REINFORCEMENT LEARNING VIA TWO-STEP ANALYSIS 2445

Xiaoyu Zhang, Wei Gao, School of Electronic and Computer Engineering, Peking University, Shenzhen, China. Peng Cheng Laboratory, Shenzhen, China., China; , ,

IVMSP-29.3: HIGH-FIDELITY PORTRAIT EDITING VIA EXPLORING DIFFERENTIABLE GUIDED SKETCHES FROM THE LATENT SPACE 2450

Chengrong Wang, Chenjie Cao, Yanwei Fu, Xiangyang Xue, Fudan University, China

IVMSP-29.4: LEARNING ADJUSTABLE IMAGE RESCALING WITH JOINT OPTIMIZATION OF PERCEPTION AND DISTORTION 2455

Zhihong Pan, Baidu Research (USA), United States of America

IVMSP-29.5: FSOINET: FEATURE-SPACE OPTIMIZATION-INSPIRED NETWORK FOR IMAGE COMPRESSIVE SENSING 2460

Wenjun Chen, Chunling Yang, Xin Yang, South China University of Technology, China

IVMSP-29.6: DISENTANGLED FEATURE-GUIDED MULTI-EXPOSURE HIGH DYNAMIC RANGE IMAGING 2465

Keuntek Lee, Yeong Il Jang, Nam Ik Cho, Seoul National University, Korea, Republic of

IVMSP-30: SECURITY AND IMAGE INTERPRETATION

IVMSP-30.1: DEFENDING AGAINST UNIVERSAL ATTACK VIA CURVATURE-AWARE CATEGORY ADVERSARIAL TRAINING 2470

Peilun Du, Xiaolong Zheng, Liang Liu, Huadong Ma, Beijing University of Posts and Telecommunications, China

IVMSP-30.2: SP ATTACK: SINGLE-PERSPECTIVE ATTACK FOR GENERATING ADVERSARIAL OMNIDIRECTIONAL IMAGES	2475
<i>Yunjian Zhang, Yanwei Liu, Pengwei Zhan, Liming Wang, Zhen Xu, Institute of Information Engineering, Chinese Academy of Sciences, China; Jinxia Liu, Zhejiang Wanli University, China</i>	
IVMSP-30.3: FEW-SHOT ONE-CLASS DOMAIN ADAPTATION BASED ON FREQUENCY FOR IRIS PRESENTATION ATTACK DETECTION	2480
<i>Yachun Li, Ying Lian, Jingjing Wang, Yuhui Chen, Chunmao Wang, Shiliang Pu, Hikvision Research Institute, China</i>	
IVMSP-30.4: PIXINWAV: RESIDUAL STEGANOGRAPHY FOR HIDING PIXELS IN AUDIO	2485
<i>Margarita Geleta, University of California, Irvine, United States of America; Cristina Punti, Kevin McGuinness, Dublin City University, Ireland; Jordi Pons, Dolby Laboratories, Spain; Cristian Canton, Xavier Giro-i-Nieto, Universitat Politècnica de Catalunya, Spain</i>	
IVMSP-30.5: A SEMI-HANDCRAFTED KEYPOINT DETECTOR WITH DISCRIMINATIVE FEATURE ENCODING	2490
<i>Yurui Xie, Ling Guan, Ryerson University, Canada</i>	
IVMSP-30.6: SAFARI FROM VISUAL SIGNALS: RECOVERING VOLUMETRIC 3D SHAPES	2495
<i>Antonio Agudo, Institut de Robotica i Informatica Industrial, CSIC-UPC, Spain</i>	
 IVMSP-31: IMAGE AND VIDEO REPRESENTATION	
IVMSP-31.1: COUPLED FEATURE LEARNING VIA STRUCTURED CONVOLUTIONAL SPARSE CODING FOR MULTIMODAL IMAGE FUSION	2500
<i>Farshad G. Veshki, Sergiy A. Vorobyov, Aalto university, Finland</i>	
IVMSP-31.2: DOMAINDESC: LEARNING LOCAL DESCRIPTORS WITH DOMAIN ADAPTATION	2505
<i>Rongtao Xu, Weiliang Meng, Xiaopeng Zhang, Institute of Automation, Chinese Academy of Sciences, Beijing, China, China; Changwei Wang, Institute of Automation, Chinese Academy of Sciences, China; Bin Fan, University of Science and Technology, China; Yuyang Zhang, Institute of Automation, Chinese Academy of Sciences, China; Shibiao Xu, Beijing University of Posts and Telecommunications, China</i>	
IVMSP-31.3: MULTI-HEAD RELU IMPLICIT NEURAL REPRESENTATION NETWORKS	2510
<i>Arya Aftab, Sharif University of Technology, Iran (Islamic Republic of); Alireza Morsali, McGill University, Canada; Shahrokh Ghaemmaghami, Department of Electrical Engineering, Sharif University of Technology, Tehran, Iran, Iran (Islamic Republic of)</i>	
IVMSP-31.4: AN EFFICIENT METHOD FOR MODEL PRUNING USING KNOWLEDGE DISTILLATION WITH FEW SAMPLES	2515
<i>Zhaojing Zhou, Zhuqing Jiang, Aidong Men, Haiying Wang, Beijing University of Posts and Telecommunications, China; Yun Zhou, National Radio and Television Administration, China</i>	
IVMSP-31.5: ADAPTIVE INTRA-GROUP AGGREGATION FOR CO-SALIENCY DETECTION	2520
<i>Guangyu Ren, Tianhong Dai, Tania Stathaki, Imperial College London, United Kingdom of Great Britain and Northern Ireland</i>	
IVMSP-31.6: NOVEL CLASS DISCOVERY: A DEPENDENCY APPROACH	2525
<i>Tanmoy Mukherjee, Nikos Deligiannis, VUB/IMEC, Belgium</i>	

IVMSP-32: APPLICATIONS IV

IVMSP-32.1: SINGLE-SHOT BALANCED DETECTOR FOR GEOSPATIAL OBJECT 2529 DETECTION

Yanfeng Liu, Qiang Li, Yuan Yuan, Qi Wang, Northwestern Polytechnical University,, China

IVMSP-32.2: REGULARIZED LATENT SPACE EXPLORATION FOR DISCRIMINATIVE 2534 FACE SUPER-RESOLUTION

Ruixin Shi, Junzheng Zhang, Shiming Ge, Institute of Information Engineering, Chinese Academy of Sciences; School of Cyber Security, University of Chinese Academy of Sciences, China; Yong Li, Institute of Information Engineering, Chinese Academy of Sciences, China

IVMSP-32.3: ENHANCING AND DISSECTING CROWD COUNTING BY SYNTHETIC 2539 DATA

Yi Hou, Yuheng Lu, Yuan Li, Huizhu Jia, Xiaodong Xie, National Engineering Laboratory for Video Technology, Peking University, China; Chengyang Li, School of Electronics Engineering and Computer Science, Peking University, Beijing, China, China; Liping Zhu, Key Lab of Petroleum Data Mining, China University of Petroleum (Beijing), China

IVMSP-32.4: MULTI-POSE VIRTUAL TRY-ON VIA SELF-ADAPTIVE FEATURE 2544 FILTERING

Chenghu Du, Feng Yu, Minghua Jiang, Xiong Wei, Tao Peng, Xinrong Hu, Wuhan Textile University, China

IVMSP-32.5: HISTOGRAM-GUIDED SEMANTIC-AWARE COLORIZATION 2549

Jie Zhang, Yi Xiao, Guo Chen, Qingping Sun, Fangqiang Xu, Hunan University, China; Chi-Sing Leung, City University of Hong Kong, Hong Kong

IVMSP-32.6: CONTENT PRESERVING SCALE SPACE NETWORK FOR FAST IMAGE 2554 RESTORATION FROM NOISY-BLURRY PAIRS

Green Rosh K S, Nikhil Krishnan, B H Pawan Prasad, Sachin Lomte, Samsung R&D Institute India, Bangalore, India

IVMSP-33: POINT CLOUD PROCESSING

IVMSP-33.1: FLOW-BASED POINT CLOUD COMPLETION NETWORK WITH 2559 ADVERSARIAL REFINEMENT

Rong Bao, Yurui Ren, Ge Li, Wei Gao, Peking University Shenzhen Graduate School, China; Shan Liu, Tencent, China

IVMSP-33.2: WEAKLY SUPERVISED POINT CLOUD UPSAMPLING VIA OPTIMAL 2564 TRANSPORT

Zezeng Li, Weimin Wang, Na Lei, Rui Wang, Dalian University of Technology, China

IVMSP-33.3: POINT CLOUD DENOISING USING NORMAL VECTOR-BASED GRAPH 2569 WAVELET SHRINKAGE

Ryosuke Watanabe, Keisuke Nonaka, Haruhisa Kato, Tatsuya Kobayashi, KDDI Research, Inc., Japan; Eduardo Pavez, Antonio Ortega, University of Southern California, United States of America

IVMSP-33.4: DYNAMIC POINT CLOUD INTERPOLATION 2574

Anique Akhtar, Zhu Li, University of Missouri Kansas City, United States of America; Geert Van der Auwera, Jianle Chen, Qualcomm Technologies Inc., United States of America

IVMSP-33.5: POINT CLOUD ATTRIBUTE COMPRESSION VIA CHROMA 2579 SUBSAMPLING

Shashank Nelamangala Sridhara, Eduardo Pavez, Antonio Ortega, University of Southern California, United States of America; Ryosuke Watanabe, Keisuke Nonaka, KDDI Research, Inc, Japan

IVMSP-33.6: RANGEINET: FAST LIDAR POINT CLOUD TEMPORAL INTERPOLATION 2584

Lili Zhao, Xuhu Lin, Wenyi Wang, Jianwen Chen, University of Electronic Science and Technology of China, China; Kai-Kuang Ma, Nanyang Technological University, Singapore

IVMSP-34: DEEP NEURAL NETWORKS FOR IMAGE AND VIDEO PROCESSING

IVMSP-34.1: MBNET: A MULTI-RESOLUTION BRANCH NETWORK FOR SEMANTIC SEGMENTATION OF ULTRA-HIGH RESOLUTION IMAGES 2589

Lianlei Shan, University of Chinese Academy of Sciences, China; Weiqiang Wang, University of Chinese Academy of Science, China

IVMSP-34.2: BSOLO: BOUNDARY-AWARE ONE-STAGE INSTANCE SEGMENTATION SOLO 2594

Yuxuan Zhang, Wei Yang, University of Science and Technology of China, China

IVMSP-34.3: CS-GRESNET: A SIMPLE AND HIGHLY EFFICIENT NETWORK FOR FACIAL EXPRESSION RECOGNITION 2599

Shaoping Jiang, Xiangmin Xu, Xiaofen Xing, Lin Wang, South China University of Technology, China; Fang Liu, Guangdong University of Finance, China

IVMSP-34.4: RCANET: ROW-COLUMN ATTENTION NETWORK FOR SEMANTIC SEGMENTATION 2604

Bingxu Lu, Qinghua Hu, Yu Wang, Tianjin University, China; Guosheng Hu, AnyVision, China

IVMSP-34.5: EXPLORING CATEGORY CONSISTENCY FOR WEAKLY SUPERVISED SEMANTIC SEGMENTATION 2609

Zhaozhi Xie, Lu Hongtao, Shanghai Jiao Tong University, China

IVMSP-34.6: VISION TRANSFORMER EQUIPPED WITH NEURAL RESIZER ON FACIAL EXPRESSION RECOGNITION TASK 2614

Hyeonbin Hwang, Soyeon Kim, Hyeon Yeo, KAIST, Korea, Republic of; Wei-Jin Park, Jiho Seo, Kyungtae Ko, ACRYL, Korea, Republic of

IVMSP-35: IMAGE SEGMENTATION

IVMSP-35.1: ISDA: POSITION-AWARE INSTANCE SEGMENTATION WITH DEFORMABLE ATTENTION 2619

Kaining Ying, Zhenhua Wang, Cong Bai, Pengfei Zhou, Zhejiang University of Technology, China

IVMSP-35.2: IMPROVING CLASS ACTIVATION MAP FOR WEAKLY SUPERVISED OBJECT LOCALIZATION 2624

Zhenfei Zhang, Ming-Ching Chang, University at Albany, United States of America; D.Bui Tien, Concordia University, Canada

IVMSP-35.3: A ROBUST OBJECT SEGMENTATION NETWORK FOR UNDERWATER SCENES 2629

Ruizhe Chen, Zhenqi Fu, Yue Huang, En Cheng, Xinghao Ding, Xiamen University, China

IVMSP-35.4: A FAST AND EFFICIENT NETWORK FOR SINGLE IMAGE SHADOW DETECTION 2634

Leiping Jie, Hong Kong Baptist University, Hong Kong; Hui Zhang, United International College, BNU-HKBU, China

IVMSP-35.5: IMPORTANCE SAMPLING CAMS FOR WEAKLY-SUPERVISED SEGMENTATION 2639

Arvi Jonnarth, Michael Felsberg, Linköping University, Sweden

IVMSP-35.6: DEEPGBASS: DEEP GUIDED BOUNDARY-AWARE SEMANTIC SEGMENTATION 2644

Qingfeng Liu, Hai Su, Mostafa El-Khamy, Kee-Bong Song, Samsung USA, United States of America

IVMSP-36: CAMERA CALIBRATION AND HUMAN POSE

IVMSP-36.1: CAMERA CALIBRATION THROUGH CAMERA PROJECTION LOSS	2649
<i>Talha Butt, Murtaza Taj, Lahore University of Management Sciences, Pakistan</i>	

IVMSP-36.2: INFERRING CAMERA INTRINSICS BASED ON SURFACES OF 2654	2654
REVOLUTION: A SINGLE IMAGE GEOMETRIC NETWORK APPROACH FOR CAMERA CALIBRATION	
<i>Christopher Walker, Yuxing Wang, Yawen Lu, Guoyu Lu, Rochester Institute of Technology, United States of America</i>	

IVMSP-36.3: TEXT2VIDEO: TEXT-DRIVEN TALKING-HEAD VIDEO SYNTHESIS 2659	2659
WITH PERSONALIZED PHONEME - POSE DICTIONARY	
<i>Sibo Zhang, Jiahong Yuan, Miao Liao, Liangjun Zhang, Baidu Research, United States of America</i>	

IVMSP-36.4: TOWARDS ACCURATE CROSS-DOMAIN IN-BED HUMAN POSE 2664	2664
ESTIMATION	
<i>Mohamed Afham, Udith Haputhanthri, Jathurshan Pradeepkumar, Mithunjha Anandakumar, Chamira Edussooriya, University of Moratuwa, Sri Lanka; Ashwin De Silva, Johns Hopkins University, United States of America</i>	

IVMSP-36.5: LEARNING MONOCULAR MESH RECOVERY OF MULTIPLE BODY 2669	2669
PARTS VIA SYNTHETICS	
<i>Yu Sun, Wenpeng Gao, Yili Fu, Harbin Institute of Technology, China; Tianyu Huang, University of Maryland, China; Qian Bao, Wu Liu, AI Research of JD.com, China</i>	

IVMSP-36.6: LIGHTPOSE: A LIGHTWEIGHT AND EFFICIENT MODEL WITH 2674	2674
TRANSFORMER FOR HUMAN POSE ESTIMATION	
<i>Xiyang Liu, Peng Li, Ding Ni, Yan Wang, Hui Xue, Alibaba Group, China</i>	

IVMSP-37: STATISTICAL AND STRUCTURAL MODELS

IVMSP-37.1: ON THE OBSERVABILITY IN VISUAL SLAM NETWORKS 2679	2679
<i>Qier An, Yuan Shen, Tsinghua University, China</i>	

IVMSP-37.2: VARIATIONAL BAYESIAN FRAMEWORK FOR ADVANCED IMAGE 2684	2684
GENERATION WITH DOMAIN-RELATED VARIABLES	
<i>Yuxiao Li, Yuan Shen, Tsinghua University, China; Santiago Mazuelas, BCAM-Basque Center for Applied Mathematics, Spain</i>	

IVMSP-37.3: THE IMPACT OF JPEG COMPRESSION ON PRIOR IMAGE NOISE 2689	2689
<i>Marina Gardella, Tina Nikoukhah, Yanhao Li, Quentin Bammey, École Normale Supérieure Paris-Saclay, France</i>	

IVMSP-37.5: ON THE USE OF COMPONENT STRUCTURAL CHARACTERISTICS FOR 2694	2694
VOXEL SEGMENTATION IN SEMICON 3D IMAGES	
<i>Tin Lay Nwe, Singh Pahwa Ramanpreet, Chang Richard, Zaw Min Oo, Jie Wang, Yiqun Li, Dongyun Lin, Prasad Shitala, Sheng Dong, Institute for Infocomm Research (I2R), Singapore</i>	

IVMSP-37.6: BLIND SOURCE SEPARATION VIA A WEAK EXCLUSION PRINCIPLE 2699	2699
<i>Zihan Zhang, Thierry Blu, The Chinese University of Hong Kong, Hong Kong</i>	

IVMSP-38: PERSON RE-IDENTIFICATION

IVMSP-38.1: GRAPH CONVOLUTION FOR RE-RANKING IN PERSON 2704	2704
RE-IDENTIFICATION	
<i>Yuqi Zhang, Qi Qian, Chong Liu, Weihua Chen, Fan Wang, Hao Li, Rong Jin, Alibaba Group, China</i>	

IVMSP-38.2: MULTI-LEVEL RELATION AWARE NETWORK FOR PERSON RE-IDENTIFICATION	2709
<i>Jing Yang, Canlong Zhang, Zhixin Li, Guangxi Normal University, China; Yanping Tang, Guilin University of Electronic Technology, China</i>	
IVMSP-38.3: PROGRESSIVE-GRANULARITY RETRIEVAL VIA HIERARCHICAL FEATURE ALIGNMENT FOR PERSON RE-IDENTIFICATION	2714
<i>Zhaopeng Dou, Zhongdao Wang, Yali Li, Shengjin Wang, Tsinghua University, China</i>	
IVMSP-38.4: OCCLUDED PERSON RE-IDENTIFICATION VIA RELATIONAL ADAPTIVE FEATURE CORRECTION LEARNING	2719
<i>Minjung Kim, MyeongAh Cho, Heansung Lee, Suhwan Cho, Sangyoun Lee, Yonsei University, Korea, Republic of</i>	
IVMSP-38.5: LEARNING SEMANTIC-ALIGNED FEATURE REPRESENTATION FOR TEXT-BASED PERSON SEARCH	2724
<i>Shiping Li, Min Cao, Min Zhang, Soochow University, China</i>	
IVMSP-38.6: TRANSFORMER-BASED PERSON SEARCH MODEL WITH SYMMETRIC ONLINE INSTANCE MATCHING	2729
<i>Xuezhi Xiang, Ning Lv, Yulong Qiao, Harbin Engineering University, China</i>	
IVMSP-39: SECURITY	
IVMSP-39.1: WASSERTRAIN: AN ADVERSARIAL TRAINING FRAMEWORK AGAINST WASSERSTEIN ADVERSARIAL ATTACKS	2734
<i>Qingye Zhao, Xin Chen, Enyi Tang, Xuandong Li, Nanjing University, China; Zhuoyu Zhao, Henan University of Economics and Law, China</i>	
IVMSP-39.2: EFFICIENT UNIVERSAL SHUFFLE ATTACK FOR VISUAL OBJECT TRACKING	2739
<i>Siao Liu, Zhaoyu Chen, Wei Li, Jiwei Zhu, Jiafeng Wang, Wenqiang Zhang, Zhongxue Gan, Fudan University, China</i>	
IVMSP-39.3: NON-RIGID TRANSFORMATION BASED ADVERSARIAL ATTACK AGAINST 3D OBJECT TRACKING	2744
<i>Riran Cheng, Nan Sang, Yinyuan Zhou, Xupeng Wang, University of Electronic Science and Technology of China, China</i>	
IVMSP-39.4: ADVERSARY DISTILLATION FOR ONE-SHOT ATTACKS ON 3D TARGET TRACKING	2749
<i>Zhengyi Wang, Xupeng Wang, Yong Liao, Jiali Yu, University of Electronic Science and Technology of China, China; Ferdous Sohel, Murdoch University, Australia; Mohammed Bennamoun, The University of Western Australia, Australia</i>	
IVMSP-39.5: ADVERFACIAL: PRIVACY-PRESERVING UNIVERSAL ADVERSARIAL PERTURBATION AGAINST FACIAL MICRO-EXPRESSION LEAKAGES	2754
<i>Yin Yin Low, Angeline Tanvy, Raphael C.W. Phan, Monash University, Malaysia; Xiaojun Chang, RMIT University, Australia</i>	
IVMSP-39.6: INTERPRETABLE IMAGE CLASSIFICATION USING SPARSE OBLIQUE DECISION TREES	2759
<i>Suryabhan Singh Hada, Miguel A. Carreira-Perpinan, University of California, Merced, United States of America</i>	
IVMSP-40: UNDERWATER IMAGING	
IVMSP-40.1: UNDERWATER IMAGE ENHANCEMENT VIA LEARNING WATER TYPE DESENSITIZED REPRESENTATIONS	2764
<i>Zhenqi Fu, Xiaopeng Lin, Wu Wang, Yue Huang, Xinghao Ding, Xiamen University, China</i>	

IVMSP-40.2: A WAVELET-BASED DUAL-STREAM NETWORK FOR UNDERWATER IMAGE ENHANCEMENT	2769
<i>Zi Yin Ma, Changjae Oh, Queen Mary University of London, United Kingdom of Great Britain and Northern Ireland</i>	
IVMSP-40.3: UNSUPERVISED AND UNTRAINED UNDERWATER IMAGE RESTORATION BASED ON PHYSICAL IMAGE FORMATION MODEL	2774
<i>Shu Chai, Zhenqi Fu, Yue Huang, Xiaotong Tu, Xinghao Ding, Xiamen University, China</i>	
IVMSP-40.4: AGCYCLEGAN: ATTENTION-GUIDED CYCLEGAN FOR SINGLE UNDERWATER IMAGE RESTORATION	2779
<i>Zhenlong Wang, Weifeng Liu, Yanjiang Wang, Baodi Liu, China University of Petroleum (East China), China</i>	
IVMSP-40.5: UNDERWATER SMALL TARGET DETECTION BASED ON DEFORMABLE CONVOLUTIONAL PYRAMID	2784
<i>Shuhan Qi, Jianjun Du, Mingyan Wu, Linlin Tang, Tao Qian, Xuan Wang, School of Computer Science and Technology, Harbin Institute of Technology Shenzhen, China; Hong Yi, Ricoh Software Research Center (Beijing), China</i>	
IVMSP-40.6: TOWARDS CONTROLLABLE AND PHYSICAL INTERPRETABLE UNDERWATER SCENE SIMULATION	2789
<i>Kaixin Chen, Lin Zhang, Ying Shen, Tongji University, China; Yicong Zhou, University of Macau, China</i>	
IVMSP-41: HYPESPECTRAL IMAGING	
IVMSP-41.1: GRAPH LEARNING BASED AUTOENCODER FOR HYPERSPECTRAL BAND SELECTION	2794
<i>Yongshan Zhang, Zhenyu Wang, Xinwei Jiang, China University of Geosciences, China; Xinxin Wang, Yicong Zhou, University of Macau, China</i>	
IVMSP-41.2: MULTITASK SPARSE NEURAL NETWORK FOR HYPERSPECTRAL IMAGE DENOISING	2799
<i>Fengchao Xiong, Jianfeng Lu, Nanjing University of Science and Technology, China; Minchao Ye, China Jiliang University, China; Jun Zhou, Griffith University, Australia; Yuntao Qian, College of Computer Science, China</i>	
IVMSP-41.3: HYPERSPECTRAL IMAGE CLASSIFICATION BASED ON CO-LEARNING THROUGH DUAL-ARCHITECTURE ENSEMBLE	2804
<i>Xiaoyue Chen, Xianghai Cao, Key Laboratory of Intelligent Perception and Image Understanding of Ministry of Education School of Artificial Intelligence, Xidian University, China</i>	
IVMSP-41.4: MATERIAL-GUIDED SIAMESE FUSION NETWORK FOR HYPERSPECTRAL OBJECT TRACKING	2809
<i>Zhuanfeng Li, Fengchao Xiong, Jianfeng Lu, Nanjing University of Science and Technology, China; Jun Zhou, Griffith University, Australia; Yuntao Qian, College of Computer Science, China</i>	
IVMSP-41.5: HYPERSPECTRAL IMAGE SUPER-RESOLUTION WITH DEEP PRIORS AND DEGRADATION MODEL INVERSION	2814
<i>Xiuheng Wang, Cédric Richard, Université Côte d’Azur, CNRS, OCA., France; Jie Chen, School of Marine Science and Technology, Northwestern Polytechnical University., China</i>	
IVMSP-41.6: GEOMETRIC LOW-RANK TENSOR APPROXIMATION FOR REMOTELY SENSED HYPERSPECTRAL AND MULTISPECTRAL IMAGERY FUSION	2819
<i>Na Liu, Wei Li, Ran Tao, Beijing Institute of Technology, China</i>	

IVMSP-42: VIDEO COMMUNICATION

IVMSP-42.1: DILATED CONVOLUTIONAL NEURAL NETWORK-BASED DEEP 2824 REFERENCE PICTURE GENERATION FOR VIDEO COMPRESSION

Haoyue Tian, Pan Gao, Nanjing University of Aeronautics and Astronautics, China; Ran Wei, Science and Technology on Electro-optic Control Laboratory, China; Manoranjan Paul, Charles Sturt University, Australia

IVMSP-42.2: RATE CONTROL FOR LEARNED VIDEO COMPRESSION 2829

Yanghao Li, Xinyao Chen, Jisheng Li, Jiangtao Wen, Tsinghua University, China; Yuxing Han, Research Institute of Tsinghua University in Shenzhen, China; Shan Liu, Xiaozhong Xu, Tencent Media Lab, China

IVMSP-42.3: GLOBAL OPTIMIZATION SOLUTION FOR DYNAMIC ADAPTIVE 2834 360-DEGREE STREAMING

Xuekai Wei, Weijia Jia, Beijing Normal University, China; Mingliang Zhou, Chongqing University, China

IVMSP-42.4: COLLABORATIVE OBJECT DETECTORS ADAPTIVE TO BANDWIDTH 2839 AND COMPUTATION

Juliano S. Assine, José C. S. Santos Filho, Eduardo Valle, Unicamp, Brazil

IVMSP-42.5: MA-NET: MULTI-SCALE ATTENTION-AWARE NETWORK FOR OPTICAL 2844 FLOW ESTIMATION

Mu Li, Baojiang Zhong, Soochow University, China; Kai-Kuang Ma, Nanyang Technological University, Singapore

IVMSP-42.6: MODELING HUMAN MEMORY IN MULTI-OBJECT TRACKING WITH 2849 TRANSFORMERS

Yizhuo Li, Cewu Lu, Shanghai Jiao Tong University, China

IFS-1: MULTIMEDIA FORENSICS I

IFS-1.1: REAL-WORLD ADVERSARIAL EXAMPLES VIA MAKEUP 2854

Chang-Sheng Lin, National Chung Hsing University, Taiwan, Province of China; Chia-Yi Hsu, Chia-Mu Yu, National Yang Ming Chiao Tung University, Taiwan, Province of China; Pin-Yu Chen, IBM Research, Taiwan, Province of China

IFS-1.2: IN PURSUIT OF PRESERVING THE FIDELITY OF ADVERSARIAL IMAGES 2859

Joseph Clements, Yingjie Lao, Clemson University, United States of America

IFS-1.3: OBJECT-ORIENTED BACKDOOR ATTACK AGAINST IMAGE CAPTIONING..... 2864

Meiling Li, Nan Zhong, Xinpeng Zhang, Zhenxing Qian, Sheng Li, Fudan University, China

IFS-1.4: TOWARDS ROBUST SPEECH-TO-TEXT ADVERSARIAL ATTACK 2869

Mohammad Esmaeilpour, Patrick Cardinal, Alessandro Lameiras Koerich, École de Technologie Supérieure, Canada

IFS-1.5: SPARSE ADVERSARIAL ATTACK FOR VIDEO VIA GRADIENT-BASED 2874 KEYFRAME SELECTION

Yixiao Xu, Xiaolei Liu, Mingyong Yin, Teng Hu, Kangyi Ding, China Academy of Engineering Physics, China

IFS-1.6: HOW SECURE ARE THE ADVERSARIAL EXAMPLES THEMSELVES?..... 2879

Hui Zeng, Southwest university of science and technology, China; Kang Deng, Anjie Peng, Southwest University of Science and Technology, China; Biwei Chen, Houghton College, United States of America

IFS-2: MULTIMEDIA FORENSICS II

IFS-2.1: EXPLORING COMPLEMENTARITY OF GLOBAL AND LOCAL 2884 SPATIOTEMPORAL INFORMATION FOR FAKE FACE VIDEO DETECTION

Xiaohui Zhao, Yang Yu, Rongrong Ni, Yao Zhao, Beijing Jiaotong University, China

IFS-2.2: PANCHROMATIC IMAGERY COPY-PASTE LOCALIZATION THROUGH DATA-DRIVEN SENSOR ATTRIBUTION	2889
<i>Edoardo Daniele Cannas, Paolo Bestagini, Stefano Tubaro, Politecnico di Milano, Italy; János Horváth, Sriram Baireddy, Edward J. Delp, Purdue University, Italy</i>	
IFS-2.3: ROBUST VIDEO HASHING BASED ON LOCAL FLUCTUATION PRESERVING FOR TRACKING DEEP FAKE VIDEOS	2894
<i>Lv Chen, Dengpan Ye, Jiaqing Huang, wuhan university, China; Yueyun Shang, South Central University for Nationalities, China</i>	
IFS-2.4: ADT: ANTI-DEEPPFAKE TRANSFORMER	2899
<i>Ping Wang, Kunlin Liu, Wenbo Zhou, Honggu Liu, Weiming Zhang, Nenghai Yu, University of Science and Technology of China, China; Hang Zhou, Simon Fraser University, Canada</i>	
IFS-2.5: EYES TELL ALL: IRREGULAR PUPIL SHAPES REVEAL GAN-GENERATED FACES	2904
<i>Hui Guo, Ming-Ching Chang, SUNY at Albany, United States of America; Shu Hu, Siwei Lyu, SUNY at Buffalo, United States of America; Xin Wang, keya medical, United States of America</i>	
IFS-2.6: EXPLAINABLE ARTIFICIAL INTELLIGENCE FOR AUTHORSHIP ATTRIBUTION ON SOCIAL MEDIA	2909
<i>Antonio Theophilo, Rafael Padilha, Fernanda A. Andaló, Anderson Rocha, University of Campinas, Brazil</i>	
IFS-3: FORENSICS AND BIOMETRICS	
IFS-3.1: DUAL-DOMAIN LOW-RANK FUSION DEEP METRIC LEARNING FOR OFF-THE-PERSON ECG BIOMETRICS	2914
<i>Guiping Zhu, Mingzhu Ma, Kuikui Wang, Gongping Yang, Shandong University, China; Yuwen Huang, Heze University, China</i>	
IFS-3.2: A ROBUST DEEP AUDIO SPLICING DETECTION METHOD VIA SINGULARITY DETECTION FEATURE	2919
<i>Kanghao Zhang, Shulin He, Jiahui Pan, Xueliang Zhang, College of Computer Science, Inner Mongolia University, China, China; Shan Liang, Shuai Nie, Haoxin Ma, Jiangyan Yi, National Laboratory of Pattern Recognition, Institute of Automation, Chinese Academy of Sciences, China</i>	
IFS-3.3: ONLINE ECG BIOMETRICS VIA HADAMARD CODE	2924
<i>Kuikui Wang, Gongping Yang, Yilong Yin, Shandong University, China; Yuwen Huang, Heze University, China; Lu Yang, Shandong Jianzhu University, China</i>	
IFS-3.4: FORENSIC ANALYSIS AND LOCALIZATION OF MULTIPLY COMPRESSED MP3 AUDIO USING TRANSFORMERS	2929
<i>Ziyue Xiang, Edward Delp, Purdue University, United States of America; Paolo Bestagini, Stefano Tubaro, Politecnico di Milano, Italy</i>	
IFS-4: SURVEILLANCE, BIOMETRICS AND SECURITY	
IFS-4.1: ADAPTIVE MATCHING STRATEGY FOR MULTI-TARGET MULTI-CAMERA TRACKING	2934
<i>Chong Liu, University of Chinese Academy of Sciences, China; Yuqi Zhang, Weihua Chen, Fan Wang, Hao Li, Alibaba Group, China; Yi-Dong Shen, Institute of Software Chinese Academy of Sciences, China</i>	
IFS-4.2: GENERALIZED FACE ANTI-SPOOFING VIA CROSS-ADVERSARIAL DISENTANGLEMENT WITH MIXING AUGMENTATION	2939
<i>Hanye Huang, Youjun Xiang, Guodong Yang, Lingling Lv, Xianfeng Li, Zichun Weng, Yuli Fu, South China University of Technology, China</i>	

IFS-4.3: FREE LUNCH FOR CROSS-DOMAIN OCCLUDED FACE RECOGNITION WITHOUT SOURCE DATA	2944
<i>Taoshan Zhang, Youjun Xiang, Xianfeng Li, Zichun Weng, Zhen Chen, Yuli Fu, South China University of Technology, China</i>	
IFS-4.4: CONEFACE: APPROXIMATE PAIRWISE LOSS FOR FACE RECOGNITION	2949
<i>Zijun Zhuang, Hongtao Lu, Shanghai Jiao Tong University, China</i>	
IFS-4.5: DEPTH-BASED ENSEMBLE LEARNING NETWORK FOR FACE ANTI-SPOOFING	2954
<i>Jie Jiang, Yunlian Sun, Nanjing University of Science and Technology, China</i>	
IFS-4.6: ARE GAN-BASED MORPHS THREATENING FACE RECOGNITION?	2959
<i>Eklavya Sarkar, Pavel Korshunov, Laurent Colbois, Sébastien Marcel, Idiap Research Institute, Switzerland</i>	
 IFS-5: PRIVACY, INFORMATION, AND NETWORK SECURITY	
IFS-5.1: PRIVACY PROTECTION IN LEARNING FAIR REPRESENTATIONS	2964
<i>Yulu Jin, Lifeng Lai, University of California, Davis, United States of America</i>	
IFS-5.2: STEALTHY BACKDOOR ATTACK WITH ADVERSARIAL TRAINING	2969
<i>Le Feng, Sheng Li, Zhenxing Qian, Xinpeng Zhang, Fudan university, China</i>	
IFS-5.3: FLDP: FLEXIBLE STRATEGY FOR LOCAL DIFFERENTIAL PRIVACY	2974
<i>Dan Zhao, Hong Chen, Suyun Zhao, Ruixuan Liu, Cuiping Li, Xiaoying Zhang, Renmin University of China, China</i>	
IFS-5.4: ENHANCING UTILITY IN THE WATCHDOG PRIVACY MECHANISM	2979
<i>Mohammad Amin Zarrabian, Australian National University, Australia; Ni Ding, University of Melbourne, Australia; Parastoo Sadeghi, University of New South Wales, Australia; Thierry Rakotoarivelo, The Commonwealth Scientific and Industrial Research Organisation, Australia</i>	
IFS-5.5: CYBER-THREAT PROPAGATION OVER NETWORK-SLICING ARCHITECTURES	2984
<i>Michele Cirillo, Mario Di Mauro, Vincenzo Matta, Giuseppe Basileo, University of Salerno, Italy</i>	
IFS-5.6: PRIVACY-AWARE COMMUNICATION OVER A WIRETAP CHANNEL WITH GENERATIVE NETWORKS	2989
<i>Ecenaz Erdemir, Pier Luigi Dragotti, Deniz Gunduz, Imperial College London, United Kingdom of Great Britain and Northern Ireland</i>	
 IFS-6: ANONYMIZATION, SECURITY, AND PRIVACY	
IFS-6.1: ENCRYPTED IMAGE VISUAL SECURITY INDEX VIA NON-LOCAL RECOGNIZABLE DEGREE EVALUATION	2994
<i>Ran Shi, Nanjing University of Science and Technology, China; Jian Xiong, Nanjing University of Posts and Telecommunications, China; Tong Qiao, Hangzhou Dianzi University, China</i>	
IFS-6.2: DEFENDING AGAINST BACKDOOR ATTACKS IN FEDERATED LEARNING WITH DIFFERENTIAL PRIVACY	2999
<i>Lu Miao, Wei Yang, Rong Hu, Liusheng Huang, University of Science and Technology of China, China; Lu Li, Yancheng Teachers University, China</i>	
IFS-6.3: SECMPNN: 3-PARTY PRIVACY-PRESERVING MOLECULAR STRUCTURE PROPERTIES INFERENCE	3004
<i>Xinying Liao, Jiaye Xue, Ximeng Liu, Fuzhou University, China; Shengxing Yu, Peking University, China; Jiangang Shu, Peng Cheng Laboratory, China</i>	

IFS-6.4: COMPRESSED DATA SHARING BASED ON INFORMATION BOTTLENECK MODEL	3009
<i>Behrooz Razeghi, Shideh Rezaeifar, Taras Holotyak, Slava Voloshynovskiy, University of Geneva, Switzerland; Sohrab Ferdowsi, HES-SO Geneva, Switzerland</i>	
IFS-6.5: RANDOMIZED SMOOTHING UNDER ATTACK: HOW GOOD IS IT IN PRACTICE?	3014
<i>Thibault Maho, Teddy Furon, Erwan Le Merrer, INRIA Rennes, France</i>	
IFS-6.6: TRAINING PRIVACY-PRESERVING VIDEO ANALYTICS PIPELINES BY SUPPRESSING FEATURES THAT REVEAL INFORMATION ABOUT PRIVATE ATTRIBUTES	3019
<i>Chau Yi Li, Andrea Cavallaro, Queen Mary University of London, United Kingdom of Great Britain and Northern Ireland</i>	
IFS-7: APPLIED CRYPTOGRAPHY AND CYBERSECURITY	
IFS-7.1: UNSUPERVISED ANOMALY DETECTION FOR CONTAINER CLOUD VIA BILSTM-BASED VARIATIONAL AUTO-ENCODER	3024
<i>Yulong Wang, Xingshu Chen, Qixu Wang, Sichuan University, China; Run Yang, Bangzhou Xin, China Academy of Engineering Physics, China</i>	
IFS-7.2: APPLYING DEEP LEARNING TO KNOWN-PLAINTEXT ATTACK ON CHAOTIC IMAGE ENCRYPTION SCHEMES	3029
<i>Fusen Wang, Jun Sang, Chunlin Huang, Bin Cai, Hong Xiang, Chongqing University, China; Nong Sang, Huazhong University of Science and Technology, China</i>	
IFS-7.3: WORDMARKOV: A NEW PASSWORD PROBABILITY MODEL OF SEMANTICS	3034
<i>Jiahong Xie, Haibo Cheng, Rong Zhu, Ping Wang, Peking University, China; Kaitai Liang, Delft University of Technology, Netherlands</i>	
IFS-7.4: EFFICIENT IDENTITY-BASED CHAMELEON HASH FOR MOBILE DEVICES	3039
<i>Cong Li, Qingni Shen, Zhikang Xie, Jisheng Dong, Yuejian Fang, Zhonghai Wu, Peking University, China</i>	
IFS-7.5: PASSTRANS: AN IMPROVED PASSWORD REUSE MODEL BASED ON TRANSFORMER	3044
<i>Xiaoxi He, Haibo Cheng, Jiahong Xie, Ping Wang, Peking University, China; Kaitai Liang, Delft University of Technology, Netherlands</i>	
IFS-8: WATERMARKING AND DATA HIDING	
IFS-8.1: FOSTERING THE ROBUSTNESS OF WHITE-BOX DEEP NEURAL NETWORK WATERMARKS BY NEURON ALIGNMENT	3049
<i>Fangqi Li, Shilin Wang, Yun Zhu, Shanghai Jiao Tong University, China</i>	
IFS-8.2: WATERMARKING IMAGES IN SELF-SUPERVISED LATENT SPACES	3054
<i>Pierre Fernandez, Alexandre Sablayrolles, Hervé Jégou, Matthijs Douze, Meta AI, France; Teddy Furon, Univ. Rennes, Inria, CNRS, IRISA, France</i>	
IFS-8.3: SPEECH PATTERN BASED BLACK-BOX MODEL WATERMARKING FOR AUTOMATIC SPEECH RECOGNITION	3059
<i>Haozhe Chen, Weiming Zhang, Kunlin Liu, Kejiang Chen, Han Fang, Nenghai Yu, University of Science and Technology of China, China</i>	
IFS-8.4: ENCRYPTION RESISTANT DEEP NEURAL NETWORK WATERMARKING	3064
<i>Guobiao Li, Sheng Li, Zhenxing Qian, Xinpeng Zhang, Fudan university, China</i>	

IFS-8.5: ATTRIBUTABLE WATERMARKING OF SPEECH GENERATIVE MODELS	3069
<i>Yongbaek Cho, Changhoon Kim, Yezhou Yang, Yi Ren, Arizona State University, United States of America</i>	
 IFS-9: STEGANOGRAPHY, STEGANALYSIS, AND DATA PRIVACY	
IFS-9.1: EXPLOITING LANGUAGE MODEL FOR EFFICIENT LINGUISTIC	3074
STEGANALYSIS	
<i>Biao Yi, Hanzhou Wu, Guorui Feng, Xinpeng Zhang, Shanghai University, China</i>	
IFS-9.2: PATCH STEGANALYSIS: A SAMPLING BASED DEFENSE AGAINST	3079
ADVERSARIAL STEGANOGRAPHY	
<i>Chuan Qin, Na Zhao, Weiming Zhang, Nenghai Yu, University of Science and Technology of China, China</i>	
IFS-9.3: AN EFFECTIVE STEGANALYSIS FOR ROBUST STEGANOGRAPHY WITH	3084
REPETITIVE JPEG COMPRESSION	
<i>Jinliu Feng, Yaofei Wang, Kejiang Chen, Weiming Zhang, Nenghai Yu, University of Science and Technology of China, China</i>	
IFS-9.4: IMAGE STEGANALYSIS WITH CONVOLUTIONAL VISION TRANSFORMER.....	3089
<i>Ge Luo, Ping Wei, Shuwen Zhu, Xinpeng Zhang, Zhenxing Qian, Sheng Li, Fudan University, China</i>	
IFS-9.5: A BRIDGE BETWEEN FEATURES AND EVIDENCE FOR BINARY	3094
ATTRIBUTE-DRIVEN PERFECT PRIVACY	
<i>Paul-Gauthier No�, Driss Matrouf, Pierre-Michel Bousquet, Jean-Fran�ois Bonastre, Avignon University, France; Andreas Nautsch, vitas.ai, Germany</i>	
IFS-9.6: PRESERVING TRAJECTORY PRIVACY IN DRIVING DATA RELEASE.....	3099
<i>Yi Xu, Chong Xiao Wang, Yang Song, Wee Peng Tay, Nanyang Technological University, Singapore</i>	
 MLSP-1: DEEP LEARNING FOR AUDIO AND MUSIC APPLICATIONS	
MLSP-1.1: DIRECT DESIGN OF BIQUAD FILTER CASCADES WITH DEEP LEARNING	3104
BY SAMPLING RANDOM POLYNOMIALS	
<i>Joseph Colonel, Christian Steinmetz, Joshua Reiss, Queen Mary University of London, United Kingdom of Great Britain and Northern Ireland; Marcus Michelen, University of Illinois at Chicago, United States of America</i>	
MLSP-1.2: AN END-TO-END DEEP LEARNING SPEECH CODING AND DENOISING	3109
STRATEGY FOR COCHLEAR IMPLANTS	
<i>Tom Gajecski, Waldo Nogueira, Hannover Medical School, Germany</i>	
MLSP-1.3: EXPLOITING HYBRID MODELS OF TENSOR-TRAIN NETWORKS FOR	3114
SPOKEN COMMAND RECOGNITION	
<i>Jun Qi, Georgia Institute of Technology, United States of America; Javier Tejedor, Universidad San Pablo-CEU, CEU Universities, Spain</i>	
MLSP-1.4: LEARNABLE WAVELET PACKET TRANSFORM FOR DATA-ADAPTED	3119
SPECTROGRAMS	
<i>Ga�tan Frusque, Olga Fink, ETH Z�rich, Switzerland</i>	
MLSP-1.5: MUSIC ENHANCEMENT VIA IMAGE TRANSLATION AND VOCODING.....	3124
<i>Nikhil Kandpal, University of North Carolina at Chapel Hill, United States of America; Oriol Nieto, Zeyu Jin, Adobe, United States of America</i>	
MLSP-1.6: PROGRESSIVE TEACHER-STUDENT TRAINING FRAMEWORK FOR	3129
MUSIC TAGGING	
<i>Rui Lu, Baigong Zheng, Jiarui Hai, Fei Tao, Zhiyao Duan, Ji Liu, Kuaishou Technology, China</i>	

MLSP-2: PATTERN RECOGNITION AND CLASSIFICATION I

MLSP-2.1: JOINT DUAL-DOMAIN MATRIX FACTORIZATION FOR ECG BIOMETRIC RECOGNITION 3134

Kuikui Wang, Gongping Yang, Yilong Yin, Shandong University, China; Yuwen Huang, Heze University, China; Lu Yang, Shandong Jianzhu University, China

MLSP-2.2: ITERATIVE SELF KNOWLEDGE DISTILLATION --- FROM POT HOLE CLASSIFICATION TO FINE-GRAINED AND COVID RECOGNITION 3139

Kuan-Chuan Peng, Mitsubishi Electric Research Laboratories (MERL), United States of America

MLSP-2.3: ATTENTION-BASED ADVERSARIAL PARTIAL DOMAIN ADAPTATION 3144

Mengzhu Wang, Xiong Peng, Wei Yu, Zhigang Luo, National University of Defense Technology, China; Shan An, Tech & Data Center JD.COM Inc, China; Xiao Luo, Peking University, China; Junyang Chen, Shenzhen University, China

MLSP-2.4: GROUP-WISE FEATURE SELECTION FOR SUPERVISED LEARNING..... 3149

Qi Xiao, Hebi Li, Jin Tian, Zhengdao Wang, Iowa State University, United States of America

MLSP-2.5: A LIGHT WEIGHT MODEL FOR VIDEO SHOT OCCLUSION DETECTION 3154

Junhua Liao, Wanbin Zhao, Yanbing Yang, Liangyin Chen, Sichuan University, China; Haihan Duan, The Chinese University of Hong Kong, China

MLSP-2.6: DETECTING BACKDOOR ATTACKS AGAINST POINT CLOUD CLASSIFIERS..... 3159

Zhen Xiang, David J. Miller, Xi Li, George Kesidis, Pennsylvania State University, United States of America; Siheng Chen, Shanghai Jiao Tong University, China

MLSP-3: SELF-SUPERVISED LEARNING FOR SPEECH AND AUDIO PROCESSING I

MLSP-3.1: CHARACTERIZING THE ADVERSARIAL VULNERABILITY OF SPEECH SELF-SUPERVISED LEARNING 3164

Haibin Wu, Hung-yi Lee, National Taiwan University, China; Bo Zheng, Xu Li, Xixin Wu, Helen Meng, The Chinese University of Hong Kong, Hong Kong

MLSP-3.2: UNIVERSAL PARALINGUISTIC SPEECH REPRESENTATIONS USING SELF-SUPERVISED CONFORMERS 3169

Joel Shor, Verily Life Sciences, United States of America; Aren Jansen, Wei Han, Daniel Park, Yu Zhang, Google, United States of America

MLSP-3.3: A NOISE-ROBUST SELF-SUPERVISED PRE-TRAINING MODEL BASED SPEECH REPRESENTATION LEARNING FOR AUTOMATIC SPEECH RECOGNITION 3174

Qiu-Shi Zhu, Jie Zhang, Zi-Qiang Zhang, Ming-Hui Wu, Xin Fang, Li-Rong Dai, University of Science and Technology of China, China

MLSP-3.4: AN ADAPTER BASED PRE-TRAINING FOR EFFICIENT AND SCALABLE SELF-SUPERVISED SPEECH REPRESENTATION LEARNING 3179

Samuel Kessler, University of Oxford, United Kingdom of Great Britain and Northern Ireland; Bethan Thomas, Salah Karout, Huawei R&D UK, United Kingdom of Great Britain and Northern Ireland

MLSP-3.5: DRVC: A FRAMEWORK OF ANY-TO-ANY VOICE CONVERSION WITH SELF-SUPERVISED LEARNING 3184

Qiqi Wang, The University of Auckland, New Zealand; Xulong Zhang, Jianzong Wang, Ning Cheng, Jing Xiao, Ping An Technology (Shenzhen) Co., Ltd., China

MLSP-3.6: CONTRASTIVE PREDICTION STRATEGIES FOR UNSUPERVISED SEGMENTATION AND CATEGORIZATION OF PHONEMES AND WORDS 3189

Santiago Cuervo, Maciej Grabias, Grzegorz Ciesielski, Paweł Rychlikowski, University of Wroclaw, Poland; Jan Chorowski, University of Wroclaw / NavAlgo, Poland; Adrian Łańcucki, NVIDIA, Poland; Ricard Marxer, Université de Toulon, Aix Marseille Univ, CNRS, LIS, France

MLSP-4: DEEP LEARNING I

MLSP-4.1: UNCERTAINTY IN DATA-DRIVEN KALMAN FILTERING FOR PARTIALLY 3194 KNOWN STATE-SPACE MODELS

Itzik Klein, University of Haifa, Israel; Guy Revach, Jonas E. Mehr, ETH Zürich, Switzerland; Nir Shlezinger, Ben-Gurion University of the Negev, Israel; Ruud J. G. van Sloun, Eindhoven University of Technology, and with Phillips Research, Netherlands; Yonina. C. Eldar, Weizmann Institute of Science, Israel

MLSP-4.2: DEEP PIECEWISE HASHING FOR EFFICIENT HAMMING SPACE 3199 RETRIEVAL

Jingzi Gu, Dayan Wu, Peng Fu, Bo Li, Weiping Wang, Institute of Information Engineering, Chinese Academy of Sciences, China

MLSP-4.3: SODA: SELF-ORGANIZING DATA AUGMENTATION IN DEEP NEURAL 3204 NETWORKS - APPLICATION TO BIOMEDICAL IMAGE SEGMENTATION TASKS

Arnaud Deleruyelle, John Klein, Cristian Versari, Université de Lille, France

MLSP-4.4: DEEP IMPULSE RESPONSES: ESTIMATING AND PARAMETERIZING 3209 FILTERS WITH DEEP NETWORKS

Alexander Richard, Peter Dodds, Vamsi Krishna Ithapu, Facebook Inc, United States of America

MLSP-4.5: JOINT TEMPORAL CONVOLUTIONAL NETWORKS AND ADVERSARIAL 3214 DISCRIMINATIVE DOMAIN ADAPTATION FOR EEG-BASED CROSS-SUBJECT EMOTION RECOGNITION

Zhipeng He, Yongshi Zhong, Jiahui Pan, South China Normal University, China

MLSP-4.6: GRADIENT VARIANCE LOSS FOR STRUCTURE-ENHANCED IMAGE 3219 SUPER-RESOLUTION

Lusine Abrahamyan, Nikos Deligiannis, Vrije Universiteit Brussel, Belgium; Anh Minh Truong, Wilfried Philips, Ghent University, Belgium

MLSP-5: PATTERN RECOGNITION AND CLASSIFICATION II

MLSP-5.1: LABEL-OCCURRENCE-BALANCED MIXUP FOR LONG-TAILED 3224 RECOGNITION

Shaoyu Zhang, Chen Chen, Silong Peng, Institute of Automation, Chinese Academy of Sciences, China; Xiujuan Zhang, Inner Mongolia Key Laboratory of Molecular Biology on Featured Plants, China

MLSP-5.2: TNTC: TWO-STREAM NETWORK WITH TRANSFORMER-BASED 3229 COMPLEMENTARITY FOR GAIT-BASED EMOTION RECOGNITION

Chuanfei Hu, Weijie Sheng, Xinde Li, Southeast University, China; Bo Dong, Zhejiang University, China

MLSP-5.3: A FREE LUNCH FROM VIT: ADAPTIVE ATTENTION MULTI-SCALE 3234 FUSION TRANSFORMER FOR FINE-GRAINED VISUAL RECOGNITION

Yuan Zhang, Jian Cao, Ling Zhang, Xiangcheng Liu, Zhiyi Wang, Feng Ling, Peking University, China; Weiqian Chen, Peking University, Alibaba Group, Inc., China

MLSP-5.4: SELF-SUPERVISED CONTRASTIVE LEARNING FOR CROSS-DOMAIN 3239 HYPERSPPECTRAL IMAGE REPRESENTATION

Hyungtae Lee, Heesung Kwon, DEVCOM Army Research Laboratory, United States of America

MLSP-5.5: GOS: A LARGE-SCALE ANNOTATED OUTDOOR SCENE SYNTHETIC 3244 DATASET

Mingye Xie, Ting Liu, Yuzhuo Fu, Shanghai Jiao Tong University, China

MLSP-5.6: OUT-OF-DISTRIBUTION AS A TARGET CLASS IN SEMI-SUPERVISED 3249 LEARNING

Antoine Tadros, Sébastien Drouyer, Rafael Grompone von Gioi, Ecole Normale Supérieure Paris-Saclay, France

MLSP-6: SELF-SUPERVISED LEARNING FOR SPEECH AND AUDIO PROCESSING II

MLSP-6.1: SELF-SUPERVISED ACOUSTIC ANOMALY DETECTION VIA 3253 CONTRASTIVE LEARNING

Hadi Hojjati, Narges Armanfard, McGill University, Canada

MLSP-6.2: DON'T SPEAK TOO FAST: THE IMPACT OF DATA BIAS ON 3258 SELF-SUPERVISED SPEECH MODELS

Yen Meng, Yi-Hui Chou, Andy T. Liu, Hung-yi Lee, National Taiwan University, Taiwan

MLSP-6.3: SELF-SUPERVISED LEARNING METHOD USING MULTIPLE SAMPLING 3263 STRATEGIES FOR GENERAL-PURPOSE AUDIO REPRESENTATION

Ibuki Kuroyanagi, Nagoya University, Japan; Tatsuya Komatsu, LINE Corporation, Japan

MLSP-6.4: SELF SUPERVISED REPRESENTATION LEARNING WITH DEEP 3268 CLUSTERING FOR ACOUSTIC UNIT DISCOVERY FROM RAW SPEECH

Varun Krishna, Sriram Ganapathy, LEAP lab, Indian Institute of Science, Bangalore, India., India

MLSP-6.5: T-NGA: TEMPORAL NETWORK GRAFTING ALGORITHM FOR LEARNING 3273 TO PROCESS SPIKING AUDIO SENSOR EVENTS

Shu Wang, Yuhuang Hu, Shih-Chii Liu, Institute of Neuroinformatics, University of Zürich and ETH Zürich, Switzerland

MLSP-6.6: CONTRASTIVE KNOWLEDGE GRAPH ATTENTION NETWORK FOR 3278 REQUEST-BASED RECIPE RECOMMENDATION

Xiyao Ma, Zheng Gao, Qian Hu, Mohamed AbdelHady, Amazon Alexa AI, United States of America

MLSP-7: DEEP LEARNING FOR IMAGE AND VIDEO PROCESSING I

MLSP-7.1: TARGETDROP: A TARGETED REGULARIZATION METHOD FOR 3283 CONVOLUTIONAL NEURAL NETWORKS

Hui Zhu, Xiaofang Zhao, Institute of Computing Technology, Chinese Academy of Sciences, China

MLSP-7.2: COARSE-TO-FINE UNSUPERVISED CHANGE DETECTION FOR REMOTE 3288 SENSING IMAGES VIA OBJECT-BASED MRF AND INCEPTION UNET

Xuan Hou, Haonan Shi, Ying Li, Northwestern Polytechnical University, China; Yunpeng Bai, Aberystwyth University, China

MLSP-7.3: COMBATING FALSE SENSE OF SECURITY: BREAKING THE DEFENSE 3293 OF ADVERSARIAL TRAINING VIA NON-GRADIENT ADVERSARIAL ATTACK

Mingyuan Fan, Wenzhong Guo, Ximeng Liu, Fuzhou University, China; Yang Liu, Xidian University, China; Gen Chen, East China Normal University, China; Shengxing Yu, Peking University, China

MLSP-7.4: DYNAMICALLY PRUNING SEGFORMER FOR EFFICIENT SEMANTIC 3298 SEGMENTATION

Haoli Bai, The Chinese University of Hong Kong, China; Hongda Mao, Dinesh Nair, Amazon.com, United States of America

MLSP-7.5: DEFORMABLE VISTR: SPATIO TEMPORAL DEFORMABLE ATTENTION 3303 FOR VIDEO INSTANCE SEGMENTATION

Sudhir Yarram, Jialian Wu, Junsong Yuan, University at Buffalo, United States of America; Pan Ji, Yi Xu, OPPO US Research Center, InnoPeak Technology Inc., United States of America

MLSP-8: PATTERN RECOGNITION AND CLASSIFICATION III

MLSP-8.1: ATTENTIONAL GATED RES2NET FOR MULTIVARIATE TIME SERIES 3308 CLASSIFICATION

Chao Yang, Xianzhi Wang, University of Technology, Sydney, Australia; Lina Yao, University of New South Wales, Australia; Guodong Long, Jing Jiang, Guandong Xu, University of Technology Sydney, Australia

MLSP-8.2: CONVEX CLUSTERING FOR AUTOCORRELATED TIME SERIES..... 3313

Max Revay, Victor Solo, University of New South Wales, Australia

MLSP-8.3: INVESTIGATING THE POTENTIAL OF AUXILIARY-CLASSIFIER GANS FOR 3318 IMAGE CLASSIFICATION IN LOW DATA REGIMES

Amil Dravid, Florian Schiffers, Yunan Wu, Oliver Cossairt, Aggelos Katsaggelos, Northwestern University, United States of America

MLSP-8.4: FEATURE AUGMENTATION LEARNING FOR FEW-SHOT PALMPRINT 3323 IMAGE RECOGNITION WITH UNCONSTRAINED ACQUISITION

Kunlei Jing, Xinman Zhang, Xi'an Jiaotong University, China; Zhiyuan Yang, Bihan Wen, Nanyang Technological University, China

MLSP-8.5: PRIME KNOWLEDGE WITH LOCAL PATTERN CONSISTENCY FOR 3328 KNOWLEDGE DISTILLATION

Qiankun Tang, Jun Wang, Zhejiang Lab, China; Xiaogang Xu, Zhejiang Gongshang University, China

MLSP-8.6: TEST-TIME DETECTION OF BACKDOOR TRIGGERS FOR POISONED 3333 DEEP NEURAL NETWORKS

Xi Li, Zhen Xiang, David Miller, George Kesidis, Pennsylvania State University, United States of America

MLSP-9: MATRIX FACTORIZATION AND COMPLETION: THEORY AND APPLICATIONS

MLSP-9.1: MULTI-VIEW DATA REPRESENTATION VIA DEEP AUTOENCODER-LIKE 3338 NONNEGATIVE MATRIX FACTORIZATION

Haonan Huang, Yihao Luo, Guoxu Zhou, Qibin Zhao, Guangdong University of Technology, China

MLSP-9.2: ON IDENTIFIABLE POLYTOPE CHARACTERIZATION FOR POLYTOPIC 3343 MATRIX FACTORIZATION

Bariscan Bozkurt, Alper Erdogan, Koc University, Turkey

MLSP-9.3: FAST LEARNING OF FAST TRANSFORMS, WITH GUARANTEES 3348

Quoc-Tung Le, Elisa Riccietti, Ecole normale supérieure de Lyon, France; Léon Zheng, valeo.ai, France; Rémi Gribonval, Inria, France

MLSP-9.4: REGRESSION ASSISTED MATRIX COMPLETION FOR 3353 RECONSTRUCTING A PROPAGATION FIELD WITH APPLICATION TO SOURCE LOCALIZATION

Hao Sun, Junting Chen, The Chinese University of Hong Kong(Shenzhen), China

MLSP-9.5: MATRIX DECOMPOSITION ON GRAPHS: A SIMPLIFIED FUNCTIONAL 3358 VIEW

Abhishek Sharma, Maks Ovsjanikov, Ecole Polytechnique, IP Paris, France

MLSP-10: SPARSITY AWARE PROCESSING

MLSP-10.1: LEARNING TO SAMPLE FOR SPARSE SIGNALS..... 3363

Satish Mulleti, Haiyang Zhang, Weizmann Institute of Science, Rehovot, Israel, Israel; Yonina C. Eldar, Weizmann Institute of Science, Israel

MLSP-10.2: MIXTURE MODEL AUTO-ENCODERS: DEEP CLUSTERING THROUGH DICTIONARY LEARNING	3368
<i>Alexander Lin, Demba Ba, Harvard University, United States of America; Andrew Song, Massachusetts Institute of Technology, United States of America</i>	
MLSP-10.3: EXPLORING THE EFFECT OF L0/L2 REGULARIZATION IN NEURAL NETWORK PRUNING USING THE LC TOOLKIT	3373
<i>Yerlan Idelbayev, Miguel Á. Carreira-Perpiñán, UC Merced, United States of America</i>	
MLSP-10.4: DICTIONARY LEARNING WITH UNIFORM SPARSE REPRESENTATIONS FOR ANOMALY DETECTION	3378
<i>Paul Irofti, Cristian Rusu, Andrei Patrascu, University of Bucharest, Faculty of Mathematics and Computer Science, Research Center for Logic, Optimization and Security (LOS), Romania</i>	
MLSP-10.5: DATA-DRIVEN SPATIALLY DEPENDENT PDE IDENTIFICATION	3383
<i>Ruixian Liu, Michael Bianco, Peter Gerstoft, Bhaskar Rao, University of California, San Diego, United States of America</i>	
MLSP-10.6: SPARSITY IMPROVES UNSUPERVISED ATTRIBUTE DISCOVERY IN STYLEGAN	3388
<i>Shusen Liu, Rushil Anirudh, Jayaraman J. Thiagarajan, Peer-Timo Bremer, Lawrence Livermore National Laboratory, United States of America</i>	
MLSP-11: DEEP LEARNING FOR IMAGE AND VIDEO PROCESSING II	
MLSP-11.1: IMAGE-TO-GRAPH TRANSFORMERS FOR CHEMICAL STRUCTURE RECOGNITION	3393
<i>Sanghyun Yoo, Ohyun Kwon, Hoshik Lee, Samsung Advanced Institute of Technology, Korea, Republic of</i>	
MLSP-11.2: A SIMPLE HYBRID FILTER PRUNING FOR EFFICIENT EDGE INFERENCE	3398
<i>Shabbeer Basha Shaik Hussain, Sheethal N Gowda, Jayachandra Dakala, DryvAmigo, Bangalore, India., India</i>	
MLSP-11.3: AN ENHANCED DEEP LEARNING APPROACH FOR TECTONIC FAULT AND FRACTURE EXTRACTION IN VERY HIGH RESOLUTION OPTICAL IMAGES	3403
<i>Bilel Kanoun, Géoazur, INRIA, France; Mohamed Abderrazak Cherif, Géoazur, LuxCarta Technology, France; Isabelle Manighetti, Géoazur, France; Yuliya Tarabalka, LuxCarta Technology, France; Josiane Zerubia, INRIA, France</i>	
MLSP-11.4: JOINT LEARNING OF FEATURE EXTRACTION AND COST AGGREGATION FOR SEMANTIC CORRESPONDENCE	3408
<i>Jiwon Kim, Youngjo Min, Mira Kim, Seungryong Kim, Korea University, Korea, Republic of</i>	
MLSP-11.5: GENERALIZED ZERO-SHOT LEARNING USING CONDITIONAL WASSERSTEIN AUTOENCODER	3413
<i>Junhan Kim, Byonghyo Shim, Seoul National University, Korea, Republic of</i>	
MLSP-11.6: MBA-RAINGAN: A MULTI-BRANCH ATTENTION GENERATIVE ADVERSARIAL NETWORK FOR MIXTURE OF RAIN REMOVAL	3418
<i>Yiyang Shen, Yidan Feng, Dong Liang, Mingqiang Wei, Nanjing University of Aeronautics and Astronautics, China; Weiming Wang, Hong Kong Metropolitan University, China; Jing Qin, The Hong Kong Polytechnic University, China; Haoran Xie, Lingnan University, China</i>	
MLSP-12: PATTERN RECOGNITION AND CLASSIFICATION IV	
MLSP-12.1: END-TO-END KEYWORD SPOTTING USING NEURAL ARCHITECTURE SEARCH AND QUANTIZATION	3423
<i>David Peter, Wolfgang Roth, Franz Pernkopf, Graz University of Technology, Austria</i>	

MLSP-12.2: SYNPOSE: A LARGE-SCALE AND DENSELY ANNOTATED SYNTHETIC DATASET FOR HUMAN POSE ESTIMATION IN CLASSROOM	3428
<i>Zefang Yu, Yangcheng Li, Yicheng Liu, Ting Liu, Yuzhuo Fu, Shanghai Jiao Tong University, China</i>	
MLSP-12.3: STPOINTGCN: SPATIAL TEMPORAL GRAPH CONVOLUTIONAL NETWORK FOR MULTIPLE PEOPLE RECOGNITION USING MILLIMETER-WAVE RADAR	3433
<i>Chunyu Wang, Peixian Gong, Lihua Zhang, Fudan University, China</i>	
MLSP-12.4: MULTIPLE TEMPORAL CONTEXT EMBEDDING NETWORKS FOR UNSUPERVISED TIME SERIES ANOMALY DETECTION	3438
<i>Hanhui Li, Xinggan Peng, Huiping Zhuang, Zhiping Lin, Nanyang Technological University, Singapore</i>	
MLSP-12.5: INTERMIX: AN INTERFERENCE-BASED DATA AUGMENTATION AND REGULARIZATION TECHNIQUE FOR AUTOMATIC DEEP SOUND CLASSIFICATION	3443
<i>Ramit Sawhney, Scaler, India; Atula Tejaswi Neerkaje, Manipal Institute of Technology, India</i>	
MLSP-12.6: CROSS-LAYER AGGREGATION WITH TRANSFORMERS FOR MULTI-LABEL IMAGE CLASSIFICATION	3448
<i>Weibo Zhang, Fuqing Zhu, Songlin Hu, 1, Institute of Information Engineering, Chinese Academy of Sciences;2,School of Cyber Security, University of Chinese Academy of Sciences, China; Jizhong Han, Tao Guo, Institute of Information Engineering, Chinese Academy of Sciences, China</i>	
MLSP-13: SELF-SUPERVISED LEARNING METHODS I	
MLSP-13.1: VISUAL REPRESENTATION LEARNING WITH SELF-SUPERVISED ATTENTION FOR LOW-LABEL HIGH-DATA REGIME	N/A
<i>Prarthana Bhattacharyya, University of Waterloo, Canada; Chenge Li, Xiaonan Zhao, István Fehérvári, Jason Sun, Amazon Inc., Canada</i>	
MLSP-13.2: TRIBYOL: TRIPLET BYOL FOR SELF-SUPERVISED REPRESENTATION LEARNING	3458
<i>Guang Li, Ren Togo, Takahiro Ogawa, Miki Haseyama, Hokkaido University, Japan</i>	
MLSP-13.3: SAGA: SELF-AUGMENTATION WITH GUIDED ATTENTION FOR REPRESENTATION LEARNING	3463
<i>Chun-Hsiao Yeh, Academia Sinica / UC Berkeley, Taiwan; Cheng-Yao Hong, Yen-Chi Hsu, Tyng-Luh Liu, Academia Sinica, Taiwan</i>	
MLSP-13.4: AN ANOMALY DETECTION METHOD BASED ON SELF-SUPERVISED LEARNING WITH SOFT LABEL ASSIGNMENT FOR DEFECT VISUAL INSPECTION	3468
<i>Chuanfei Hu, Yongxiong Wang, University of Shanghai for Science and Technology, China</i>	
MLSP-13.5: CONTRASTIVE PREDICTIVE CODING FOR ANOMALY DETECTION OF FETAL HEALTH FROM THE CARDIOTOCOGRAM	3473
<i>Ivar R. de Vries, Iris A.M. Huijben, Ruud J.G. van Sloun, Rik Vullings, Eindhoven University of Technology, Netherlands; René D. Kok, Nemo Healthcare BV, Netherlands</i>	
MLSP-13.6: GRAPH FINE-GRAINED CONTRASTIVE REPRESENTATION LEARNING	3478
<i>Hui Tang, Xun Liang, Yuhui Guo, Xiangping Zheng, Bo Wu, Renmin University of China, China</i>	
MLSP-14: DEEP LEARNING II	
MLSP-14.1: POSITION-INVARIANT ADVERSARIAL ATTACKS ON NEURAL MODULATION RECOGNITION	3483
<i>Zhen Yu, Yifeng Xiong, Kun He, Huazhong University of Science and Technology, China; Shao Huang, Yaodong Zhao, Jie Gu, Science and Technology on Electronic Information Control Laboratory, China</i>	

MLSP-14.2: USING A SINGLE INPUT TO FORECAST HUMAN ACTION KEYSTATES IN EVERYDAY PICK AND PLACE ACTIONS	3488
<i>Haziq Razali, Yiannis Demiris, Imperial College London, United Kingdom of Great Britain and Northern Ireland</i>	
MLSP-14.3: ADVERSARIAL ROBUSTNESS BY DESIGN THROUGH ANALOG COMPUTING AND SYNTHETIC GRADIENTS	3493
<i>Alessandro Cappelli, Iacopo Poli, LightOn, France; Ruben Ohana, Julien Launay, LightOn/ Ecole Normale Supérieure, France; Laurent Meunier, Université Dauphine/ Facebook AI Research, France; Florent Krzakala, LightOn/ EPFL, France</i>	
MLSP-14.4: DIFFERENTIABLE PROGRAMMING A LA MOREAU	3498
<i>Vincent Roulet, Zaid Harchaoui, University of Washington, United States of America</i>	
MLSP-14.5: DATA AGNOSTIC FILTER GATING FOR EFFICIENT DEEP NETWORKS	3503
<i>Hongyan Xu, Arcot Sowmya, the University of New South Wales, Australia; Xiu Su, Chang Xu, the University of Sydney, Australia; Shan You, Tao Huang, Fei Wang, Chen Qian, SenseTime, China; Changshui Zhang, Tsinghua University, China; Dadong Wang, Data61, The Commonwealth Scientific and Industrial Research Organisation (CSIRO), Australia</i>	
MLSP-15: PATTERN RECOGNITION AND CLASSIFICATION V	
MLSP-15.1: NEAREST SUBSPACE SEARCH IN THE SIGNED CUMULATIVE DISTRIBUTION TRANSFORM SPACE FOR 1D SIGNAL CLASSIFICATION	3508
<i>Abu Hasnat Mohammad Rubaiyat, Mohammad Shifat-E-Rabbi, Yan Zhuang, Shiyong Li, Gustavo Rohde, University of Virginia, United States of America</i>	
MLSP-15.2: ENERGY ALIGNMENT FOR BIAS RECTIFICATION IN CLASS INCREMENTAL LEARNING	3513
<i>Bowen Zhao, Xi Xiao, Shutao Xia, Tsinghua University, China; Chen Chen, Qi Ju, Tencent, China</i>	
MLSP-15.3: A TWO-STAGE CONTRASTIVE LEARNING FRAMEWORK FOR IMBALANCED AERIAL SCENE RECOGNITION	3518
<i>Lexing Huang, Senlin Cai, Yihong Zhuang, Changxing Jing, Yue Huang, Xiaotong Tu, Xinghao Ding, Xiamen University, China</i>	
MLSP-15.4: A MAXIMAL CORRELATION APPROACH TO IMPOSING FAIRNESS IN MACHINE LEARNING	3523
<i>Joshua Lee, Yuheng Bu, Gregory Wornell, MIT, United States of America; Prasanna Sattigeri, Rameswar Panda, Leonid Karlinsky, Rogerio Feris, IBM Research, United States of America</i>	
MLSP-15.5: BOUNDARY-AWARE BIAS LOSS FOR TRANSFORMER-BASED AERIAL IMAGE SEGMENTATION MODEL	3528
<i>Yan Zhang, Xue Jiang, Bo Hu, Xinbo Gao, Chongqing University of Posts and Telecommunications, China; Siqi Liu, Peking University, China</i>	
MLSP-15.6: INVESTIGATING ROBUSTNESS OF BIOLOGICAL VS. BACKPROP BASED LEARNING	3533
<i>Yanpeng Zhou, Maosen Wang, Ponnuthurai Nagarathnam Suganthan, Nanyang Technological University, Singapore; Manas Gupta, Arulmurugan Ambikapathi, Ramasamy Savitha, A*STAR, Singapore</i>	
MLSP-16: SELF-SUPERVISED LEARNING METHODS II	
MLSP-16.1: SEMI-SUPERVISED GAUSSIAN MIXTURE VARIATIONAL AUTOENCODER FOR PULSE SHAPE DISCRIMINATION	3538
<i>Abdullah Abdulaziz, Yoann Altmann, Stephen McLaughlin, Heriot-Watt University, United Kingdom of Great Britain and Northern Ireland; Jianxin Zhou, Angela Di Fulvio, University of Illinois, United States of America</i>	

MLSP-16.2: HOW NEURAL PROCESSES IMPROVE GRAPH LINK PREDICTION	3543
<i>Huidong Liang, Junbin Gao, The University of Sydney, Australia</i>	
MLSP-16.3: UNCERTAINTY ESTIMATION WITH A VAE-CLASSIFIER HYBRID MODEL	3548
<i>Shuyu Lin, Niki Trigoni, Stephen Roberts, University of Oxford, United Kingdom of Great Britain and Northern Ireland; Ronald Clark, Imperial College London, United Kingdom of Great Britain and Northern Ireland</i>	
MLSP-16.4: CONTEXT-AWARE GRAPH-BASED SELF-SUPERVISED LEARNING OF WHOLE SLIDE IMAGES	3553
<i>Milan Aryal, Nasim Yahyasoltani, Marquette University, United States of America</i>	
MLSP-16.5: CONTRASTIVE SENSOR TRANSFORMER FOR PREDICTIVE MAINTENANCE OF INDUSTRIAL ASSETS	3558
<i>Zaharah Bukhsh, Eindhoven University of Technology, Netherlands</i>	
MLSP-16.6: IMPROVING ANOMALY DETECTION WITH A SELF-SUPERVISED TASK BASED ON GENERATIVE ADVERSARIAL NETWORK	3563
<i>Heyan Chai, Weijun Su, Siyu Tang, Harbin Institute of Technology (Shenzhen), China; Ye Ding, Dongguan University of Technology, China; Binxing Fang, Qing Liao, Harbin Institute of Technology (Shenzhen), China</i>	
MLSP-17: GRAPH NEURAL NETWORKS	
MLSP-17.1: STGAT-MAD : SPATIAL-TEMPORAL GRAPH ATTENTION NETWORK FOR MULTIVARIATE TIME SERIES ANOMALY DETECTION	3568
<i>Jun Zhan, Siqi Wang, Chengkun Wu, Detian Zeng, National University of Defense Technology, China; Xiandong Ma, Lancaster University, United Kingdom of Great Britain and Northern Ireland; Canqun Yang, National Supercomputing Center of Tianjin, China; Shilin Wang, Beijing Goldwind Smart Energy Technology Co., Ltd., China</i>	
MLSP-17.2: DUAL GRAPH CROSS-DOMAIN FEW-SHOT LEARNING FOR HYPERSPECTRAL IMAGE CLASSIFICATION	3573
<i>Yuxiang Zhang, Wei Li, Mengmeng Zhang, Ran Tao, Beijing Institute of Technology, China</i>	
MLSP-17.3: PERSONALIZED PAGERANK GRAPH ATTENTION NETWORKS	3578
<i>Julie Choi, Amazon, United States of America</i>	
MLSP-17.4: MULTI-RELATION MESSAGE PASSING FOR MULTI-LABEL TEXT CLASSIFICATION	3583
<i>Muberra Ozmen, Mark Coates, McGill University, Canada; Hao Zhang, Hong Kong University of Science and Technology, Hong Kong; Pengyun Wang, Huawei Noah's Ark Lab, China</i>	
MLSP-17.5: ADAPTIVE ATTENTION GRAPH CAPSULE NETWORK.....	3588
<i>Xiangping Zheng, Xun Liang, Bo Wu, Yuhui Guo, Hui Tang, Renmin University of China, China</i>	
MLSP-17.6: GRAPH CONVOLUTIONAL NETWORKS WITH AUTOENCODER-BASED COMPRESSION AND MULTI-LAYER GRAPH LEARNING	3593
<i>Lorenzo Giusti, Claudio Battiloro, Paolo Di Lorenzo, Sergio Barbarossa, Sapienza University of Rome, Italy</i>	
MLSP-18: MACHINE LEARNING FOR AUDIO, ACOUSTICS, AND SPEECH PROCESSING	
MLSP-18.1: DEEP AUGMENTED MUSIC ALGORITHM FOR DATA-DRIVEN DOA ESTIMATION	3598
<i>Julian P. Merkofer, Guy Revach, ETH Zurich, Switzerland; Nir Shlezinger, Ben-Gurion University of the Negev, Israel; Ruud J. G. van Sloun, Eindhoven University of Technology, Netherlands</i>	

MLSP-18.2: CONVMIXER: FEATURE INTERACTIVE CONVOLUTION WITH CURRICULUM LEARNING FOR SMALL FOOTPRINT AND NOISY FAR-FIELD KEYWORD SPOTTING	3603
<i>Dianwen Ng, Yunqi Chen, Eng Siong Chng, Nanyang Technological University, Singapore; Biao Tian, Qiang Fu, Alibaba Group, China</i>	
MLSP-18.3: SEMI-SUPERVISED SOURCE LOCALIZATION WITH RESIDUAL PHYSICAL LEARNING	3608
<i>Michael Bianco, Peter Gerstoft, University of California San Diego, United States of America</i>	
MLSP-18.4: AUTOMATED PROSODY CLASSIFICATION FOR ORAL READING FLUENCY WITH QUADRATIC KAPPA LOSS AND ATTENTIVE X-VECTORS	3613
<i>George Sammit, Zhongjie Wu, Yihao Wang, Zhongdi Wu, Akihito Kamata, Eric C. Larson, Southern Methodist University, United States of America; Joseph Nese, University of Oregon, United States of America</i>	
MLSP-18.5: SEED: SOUND EVENT EARLY DETECTION VIA EVIDENTIAL UNCERTAINTY	3618
<i>Xujiang Zhao, Feng Chen, University of Texas at Dallas, United States of America; Xuchao Zhang, Wei Cheng, Wenchao Yu, Yuncong Chen, Haifeng Chen, NEC Laboratories America, United States of America</i>	
MLSP-18.6: RANK-BASED LOSS FOR LEARNING HIERARCHICAL REPRESENTATIONS	3623
<i>Ines Nolasco, Queen Mary university of London, United Kingdom of Great Britain and Northern Ireland; Dan Stowell, Tilburg University, Netherlands</i>	
MLSP-19: TENSOR-BASED SIGNAL PROCESSING	
MLSP-19.1: ON THE RELAXATION OF ORTHOGONAL TENSOR RANK AND ITS NONCONVEX RIEMANNIAN OPTIMIZATION FOR TENSOR COMPLETION	3628
<i>Keisuke Ozawa, DENSO IT Laboratory, Japan</i>	
MLSP-19.2: ROBUST HIGH-ORDER TENSOR RECOVERY VIA NONCONVEX LOW-RANK APPROXIMATION	3633
<i>Wenjin Qin, Jianjun Wang, Southwest University, School of Mathematic and Statistics, China; Hailin Wang, Xi'an Jiaotong University, School of Mathematics and Statistics, China; Weijun Ma, Ningxia University, School of Information Engineering, China; , ,</i>	
MLSP-19.3: VARIATIONAL BAYESIAN TENSOR NETWORKS WITH STRUCTURED POSTERIOBS	3638
<i>Kriton Konstantinidis, Yao Xu, Danilo Mandic, IMPERIAL COLLEGE LONDON, United Kingdom of Great Britain and Northern Ireland; Qibin Zhao, RIKEN CENTER FOR ADVANCED INTELLIGENCE PROJECT, Japan</i>	
MLSP-19.4: LOW-RANK PHASE RETRIEVAL WITH STRUCTURED TENSOR MODELS	3643
<i>Soo Min Kwon, Xin Li, Anand Sarwate, Rutgers, The State University of New Jersey, United States of America</i>	
MLSP-19.5: HOQRI: HIGHER-ORDER QR ITERATION FOR SCALABLE TUCKER DECOMPOSITION	3648
<i>Yuchen Sun, Kejun Huang, University of Florida, United States of America</i>	
MLSP-19.6: A MULTI-RESOLUTION LOW-RANK TENSOR DECOMPOSITION	3653
<i>Sergio Rozada Doval, Antonio G. Marques, King Juan Carlos University, Spain</i>	
MLSP-20: DATA SCIENCE INITIATIVE I	
MLSP-20.1: EXPLORING HETEROGENEOUS CHARACTERISTICS OF LAYERS IN ASR MODELS FOR MORE EFFICIENT TRAINING	3658
<i>Lillian Zhou, Dhruv Guliani, Andreas Kabel, Giovanni Motta, Françoise Beaufays, Google LLC, United States of America</i>	

MLSP-20.2: PROBABILISTIC FINE-GRAINED URBAN FLOW INFERENCE WITH NORMALIZING FLOWS	3663
<i>Ting Zhong, Haoyang Yu, Rongfan Li, Xovee Xu, Xucheng Luo, Fan Zhou, University of Electronic Science and Technology of China, China</i>	
MLSP-20.3: ATTENTION-BASED DUAL-STREAM VISION TRANSFORMER FOR RADAR GAIT RECOGNITION	3668
<i>Shiliang Chen, Wentao He, Jianfeng Ren, University of Nottingham Ningbo China, China; Xudong Jiang, Nanyang Technological University, Singapore</i>	
MLSP-20.4: DEEP-MLE: FUSION BETWEEN A NEURAL NETWORK AND MLE FOR A SINGLE SNAPSHOT DOA ESTIMATION	3673
<i>Marcio L. Lima de Oliveira, University of Twente, Netherlands; Marco Bekooij, NXP Semiconductors, Netherlands</i>	
MLSP-20.5: SELECTIVE MUTUAL LEARNING: AN EFFICIENT APPROACH FOR SINGLE CHANNEL SPEECH SEPARATION	3678
<i>MinhTan Ha, Duc-Quang Vu, Chung-Ting Lee, Yung-Hui Li, Jia-Ching Wang, National Central University, Taiwan</i>	
MLSP-20.6: DETECTION OF COVID-19 FROM JOINT TIME AND FREQUENCY ANALYSIS OF SPEECH, BREATHING AND COUGH AUDIO	3683
<i>John Harvill, Moitreyia Chatterjee, Mark Hasegawa-Johnson, Narendra Ahuja, University of Illinois at Urbana-Champaign, United States of America; Yash Wani, Mustafa Alam, David Beiser, University of Chicago, United States of America; David Chestek, University of Illinois at Chicago, United States of America</i>	
MLSP-21: DEEP LEARNING FOR SPEECH AND AUDIO PROCESSING	
MLSP-21.1: SELF-CRITICAL SEQUENCE TRAINING FOR AUTOMATIC SPEECH RECOGNITION	3688
<i>Chen Chen, Yuchen Hu, Nana Hou, Heqing Zou, Eng Siong Chng, Nanyang Technological University, Singapore; Xiaofeng Qi, Hachibot, China</i>	
MLSP-21.2: FASTAUDIO: A LEARNABLE AUDIO FRONT-END FOR SPOOF SPEECH DETECTION	3693
<i>Quchen Fu, Zhongwei Teng, Jules White, Douglas Schmidt, Vanderbilt University, United States of America; Maria Powell, Vanderbilt University Medical Center, United States of America</i>	
MLSP-21.3: COMPLEX IIR-AWARE TRAINING FOR VOICE ACTIVITY DETECTION USING ATTENTION MODEL	3698
<i>Yifei Zhao, Yazid Attabi, Benoit Champagne, McGill University, Canada; Wei-Ping Zhu, Concordia University, Canada</i>	
MLSP-21.4: LEARNING CONTINUOUS REPRESENTATION OF AUDIO FOR ARBITRARY SCALE SUPER RESOLUTION	3703
<i>Jaechang Kim, Yunjoo Lee, Jungseul Ok, Pohang University of Science and Technology, Korea, Republic of; Seunghoon Hong, Korea Advanced Institute of Science and Technology, Korea, Republic of</i>	
MLSP-21.5: AN INVESTIGATION OF THE EFFECTIVENESS OF PHASE FOR AUDIO CLASSIFICATION	3708
<i>Shunsuke Hidaka, Kohei Wakamiya, Tokihiko Kaburagi, Kyushu University, Japan</i>	
MLSP-21.6: STUDY OF POSITIONAL ENCODING APPROACHES FOR AUDIO SPECTROGRAM TRANSFORMERS	3713
<i>Leonardo Pepino, Pablo Riera, Luciana Ferrer, University of Buenos Aires, Argentina</i>	
MLSP-22: MACHINE LEARNING FOR IMAGE AND VIDEO PROCESSING	
MLSP-22.1: FEW-SHOT OBJECT DETECTION WITH LOCAL CORRESPONDENCE RPN AND ATTENTIVE HEAD	3718
<i>Jian Han, Yali Li, Shengjin Wang, tsinghua.edu.cn, China</i>	

MLSP-22.2: NATURAL-LOOKING ADVERSARIAL EXAMPLES FROM FREEHAND SKETCHES	3723
<i>Hak Gu Kim, Chung-Ang University, Korea, Republic of; Davide Nanni, Sabine Süssstrunk, École Polytechnique Fédérale de Lausanne, Switzerland</i>	
MLSP-22.3: VIDEO ANOMALY DETECTION VIA PREDICTION NETWORK WITH ENHANCED SPATIO-TEMPORAL MEMORY EXCHANGE	3728
<i>Guodong Shen, Yuqi Ouyang, Victor Sanchez, University of Warwick, United Kingdom of Great Britain and Northern Ireland</i>	
MLSP-22.4: SIGNAL COMPRESSION VIA NEURAL IMPLICIT REPRESENTATIONS	3733
<i>Francesca Pistilli, Diego Valsesia, Giulia Fracastoro, Enrico Magli, Politecnico di Torino, Italy</i>	
MLSP-22.5: HYBRID WEIGHTING LOSS FOR PRECIPITATION NOWCASTING FROM RADAR IMAGES	3738
<i>Yuan Cao, Hongming Shan, Fudan University, China; Lei Chen, Leiming Ma, Shanghai Central Meteorological Observatory, China; Danchen Zhang, University of Pittsburgh, China</i>	
MLSP-22.6: ADVERSARIAL LEARNING ENHANCEMENT FOR 3D HUMAN POSE AND SHAPE ESTIMATION	3743
<i>Yidian Sun, Jiwei Zhang, Wendong Wang, Beijing University of Posts and Telecommunications, China</i>	
MLSP-23: TRANSFER LEARNING I	
MLSP-23.1: DOMAIN GENERALIZED FEW-SHOT IMAGE CLASSIFICATION VIA META REGULARIZATION NETWORK	3748
<i>Min Zhang, Zhejiang University, China; Siteng Huang, Donglin Wang, Westlake University, China</i>	
MLSP-23.2: GENERATION FOR UNSUPERVISED DOMAIN ADAPTATION: A GAN-BASED APPROACH FOR OBJECT CLASSIFICATION WITH 3D POINT CLOUD DATA	3753
<i>Junxuan Huang, Junsong Yuan, Chunming Qiao, University at buffalo, United States of America</i>	
MLSP-23.3: EXPLORING TRANSFERABILITY MEASURES AND DOMAIN SELECTION IN CROSS-DOMAIN SLOT FILLING	3758
<i>Xin-Chun Li, Yan-Jia Wang, Le Gan, De-Chuan Zhan, Nanjing University, China</i>	
MLSP-23.4: MAXIMUM BATCH FROBENIUS NORM FOR MULTI-DOMAIN TEXT CLASSIFICATION	3763
<i>Yuan Wu, Ahmed El-Roby, Carleton University, Canada; Diana Inkpen, University of Ottawa, Canada</i>	
MLSP-23.5: JOINT GLOBAL-LOCAL ALIGNMENT FOR DOMAIN ADAPTIVE SEMANTIC SEGMENTATION	3768
<i>Sudhir Yarram, Yuan Junsong, Chunming Qiao, University at Buffalo, United States of America; Ming Yang, Horizon Robotics, United States of America</i>	
MLSP-23.6: CATEGORY-ADAPTIVE DOMAIN ADAPTATION FOR SEMANTIC SEGMENTATION	3773
<i>Zhiming Wang, Yantian Luo, Danlan Huang, Ning Ge, Jianhua Lu, Tsinghua University, China</i>	
MLSP-24: DEEP GENERATIVE MODEL	
MLSP-24.1: SIMPLER IS BETTER: SPECTRAL REGULARIZATION AND UP-SAMPLING TECHNIQUES FOR VARIATIONAL AUTOENCODERS	3778
<i>Sara Björk, UiT The Arctic University of Norway, Norway; Jonas Nordhaug Myhre, Thomas Haugland Johansen, NORCE Norwegian Research Centre, Norway</i>	

MLSP-24.2: AUGMENTING MOLECULAR DEEP GENERATIVE MODELS WITH TOPOLOGICAL DATA ANALYSIS REPRESENTATIONS	3783
<i>Yair Schiff, Vijil Chenthamarakshan, Samuel Hoffman, Karthikeyan Natesan Ramamurthy, Payel Das, IBM, United States of America</i>	
MLSP-24.3: STYLEGAN-INDUCED DATA-DRIVEN REGULARIZATION FOR INVERSE PROBLEMS	3788
<i>Arthur Conmy, Subhadip Mukherjee, Carola-Bibiane Schönlieb, University of Cambridge, United Kingdom of Great Britain and Northern Ireland</i>	
MLSP-24.4: A CLOSER LOOK AT AUTOENCODERS FOR UNSUPERVISED ANOMALY DETECTION	3793
<i>Oyebade Oyedotun, Djamila Aouada, University of Luxembourg, Luxembourg</i>	
MLSP-24.5: NFT-K: NON-FUNGIBLE TANGENT KERNELS	3798
<i>Sina Alemohammad, Hossein Babaei, CJ Barberan, Naiming Liu, Lorenzo Luzi, Blake Mason, Richard Baraniuk, Rice University, United States of America</i>	
 MLSP-25: MACHINE LEARNING FOR EARTH SCIENCES AND MATERIALS	
MLSP-25.1: FDSNET: AN ACCURATE REAL-TIME SURFACE DEFECT SEGMENTATION NETWORK	3803
<i>Jian Zhang, Runwei Ding, Miaoju Ban, Tianyu Guo, Key Laboratory of Machine Perception, Peking University, Shenzhen Graduate School, China</i>	
MLSP-25.2: PATH SIGNATURES FOR NON-INTRUSIVE LOAD MONITORING	3808
<i>Paul Moore, Theodor-Mihai Iliant, Filip-Alexandru Ion, Terry Lyons, University of Oxford, United Kingdom of Great Britain and Northern Ireland; Yue Wu, University of Strathclyde, United Kingdom of Great Britain and Northern Ireland</i>	
MLSP-25.3: DATA-DRIVEN APPROACH FOR THE FLOQUET PROPAGATOR INVERSE PROBLEM SOLUTION	3813
<i>Alexander Hvatov, ITMO University, Russian Federation</i>	
MLSP-25.4: CHUNKFUSION: A LEARNING-BASED RGB-D 3D RECONSTRUCTION FRAMEWORK VIA CHUNK-WISE INTEGRATION	3818
<i>Chaozheng Guo, Lin Zhang, Ying Shen, Tongji University, China; Yicong Zhou, University of Macau, China</i>	
MLSP-25.5: CLOSING THE SIM-TO-REAL GAP IN GUIDED WAVE DAMAGE DETECTION WITH ADVERSARIAL TRAINING OF VARIATIONAL AUTO-ENCODERS	3823
<i>Ishan Khurjekar, Joel Harley, University of Florida, United States of America</i>	
MLSP-25.6: DEEP LEARNING ON THE SPHERE FOR MULTI-MODEL ENSEMBLING OF SIGNIFICANT WAVE HEIGHT	3828
<i>Andrea Littardi, Offshore Monitoring Ltd., Cyprus. The Cyprus Institute, Cyprus, Italy; Anders Hildeman, Chalmers University of Technology, Sweden., Sweden; Mihalis A. Nicolaou, The Cyprus Institute, Cyprus., Cyprus</i>	
 MLSP-26: TRANSFER LEARNING II	
MLSP-26.1: LOCAL AND GLOBAL ALIGNMENTS FOR GENERALIZABLE SENSOR-BASED HUMAN ACTIVITY RECOGNITION	3833
<i>Wang Lu, Jindong Wang, Yiqiang Chen, Institute of Computing Technology, Chinese Academy of Sciences, China</i>	
MLSP-26.2: STUDY ON TIME-OF-FLIGHT ESTIMATION IN ULTRASONIC WELL LOGGING TOOL: MODEL-DRIVEN TRANSFER LEARNING	3838
<i>Wei Zhang, Zhipeng Li, Yiduo Guo, Ao Qiu, Yanjun Li, Yibing Shi, School of Automation Engineering, University of Electronic Science and Technology of China, China</i>	

MLSP-26.3: SIMULATION-AND-MINING: TOWARDS ACCURATE SOURCE-FREE UNSUPERVISED DOMAIN ADAPTIVE OBJECT DETECTION	3843
<i>Peng Yuan, Weijie Chen, Shicai Yang, Yunyi Xuan, Di Xie, Shiliang Pu, Hikvision Research Institute, China; Yueting Zhuang, Zhejiang University, China; , ,</i>	
MLSP-26.4: TARGET-AWARE AUTO-AUGMENTATION FOR UNSUPERVISED DOMAIN ADAPTIVE OBJECT DETECTION	3848
<i>Zhaoyang Li, Long Zhao, Weijie Chen, Shicai Yang, Di Xie, Shiliang Pu, Hikvision Research Institute, China</i>	
MLSP-26.5: SELF-ENSEMBLE VARIANCE REGULARIZATION FOR DOMAIN ADAPTATION	3853
<i>Xinyi Liu, Tao Dai, Shu-Tao Xia, Yong Jiang, Tsinghua University, China</i>	
MLSP-26.6: TRANSDUCTIVE CLIP WITH CLASS-CONDITIONAL CONTRASTIVE LEARNING	3858
<i>Junchu Huang, South China University of Technology, China; Weijie Chen, Shicai Yang, Di Xie, Shiliang Pu, Hikvision Research Institute, China; Yueting Zhuang, Zhejiang University, China</i>	
MLSP-27: SUBSPACE AND MANIFOLD LEARNING	
MLSP-27.1: CONTROLLING THE FRÉCHET VARIANCE IMPROVES BATCH NORMALIZATION ON THE SYMMETRIC POSITIVE DEFINITE MANIFOLD	3863
<i>Reinmar Kobler, Jun-ichiro Hirayama, Motoaki Kawanabe, RIKEN, Japan</i>	
MLSP-27.2: SUBSPACE CLUSTERING USING UNSUPERVISED DATA AUGMENTATION	3868
<i>Maryam Abdolali, Nicolas Gillis, University of Mons, Belgium</i>	
MLSP-27.3: PRIVATE LEARNING VIA KNOWLEDGE TRANSFER WITH HIGH-DIMENSIONAL TARGETS	3873
<i>Dominik Fay, Tobias Oechtering, KTH Royal Institute of Technology, Sweden; Jens Sjölund, Uppsala University, Sweden</i>	
MLSP-27.4: DEEP DETERMINISTIC INDEPENDENT COMPONENT ANALYSIS FOR HYPERSPECTRAL UNMIXING	3878
<i>Hongming Li, Jose Principe, University of Florida, United States of America; Shujian Yu, The Arctic University of Norway, Norway</i>	
MLSP-27.5: LABEL-AWARE RANKED LOSS FOR ROBUST PEOPLE COUNTING USING AUTOMOTIVE IN-CABIN RADAR	3883
<i>Lorenzo Servadei, Huawei Sun, Julius Ott, Daniela Sánchez Lopera, Infineon Technologies AG / Technical University of Munich, Germany; Michael Stephan, Thomas Stadelmayer, Infineon Technologies AG / Friedrich-Alexander-University of Erlangen Nuremberg, Germany; Souvik Hazra, Avik Santra, Infineon Technologies AG, Germany; Robert Wille, Johannes Kepler University Linz, Germany</i>	
MLSP-27.6: DEEPHULL: FAST CONVEX HULL APPROXIMATION IN HIGH DIMENSIONS	3888
<i>Randall Balestriero, Facebook AI Research, United States of America; Zichao Wang, Richard Baraniuk, Rice University, United States of America</i>	
MLSP-28: MACHINE LEARNING FOR SOCIAL SCIENCES	
MLSP-28.1: NEIGHBOR-AUGMENTED TRANSFORMER-BASED EMBEDDING FOR RETRIEVAL	3893
<i>Jihai Zhang, Fangquan Lin, Wei Jiang, Cheng Yang, Alibaba Group, China; Gaoge Liu, Columbia University, United States of America</i>	

MLSP-28.2: SENTIMENT-AWARE DISTILLATION FOR BITCOIN TREND FORECASTING UNDER PARTIAL OBSERVABILITY	3898
<i>Georgios Panagiotatos, Nikolaos Passalis, Avraam Tsantekidis, Anastasios Tefas, Aristotle University of Thessaloniki, Greece</i>	
MLSP-28.3: ROBUST NONPARAMETRIC DISTRIBUTION FORECAST WITH BACKTEST-BASED BOOTSTRAP AND ADAPTIVE RESIDUAL SELECTION	3903
<i>Longshaokan Wang, Mina Georgieva, Paulo Machado, Abinaya Ulagappa, Safwan Ahmed, Yan Lu, Arjun Bakshi, Farhad Ghassemi, Amazon, United States of America; Lingda Wang, University of Illinois at Urbana-Champaign, United States of America</i>	
MLSP-28.4: VARIATIONAL BAYESIAN GRAPH CONVOLUTIONAL NETWORK FOR ROBUST COLLABORATIVE FILTERING	3908
<i>Nozomu Onodera, Keisuke Maeda, Takahiro Ogawa, Miki Haseyama, Hokkaido University, Japan</i>	
MLSP-28.5: FINT: FIELD-AWARE INTERACTION NEURAL NETWORK FOR CLICK-THROUGH RATE PREDICTION	3913
<i>Zhishan Zhao, Guohui Liu, iQIYI Inc., China; Sen Yang, Dawei Feng, Kele Xu, National University of Defense Technology, China</i>	
MLSP-28.6: MAKING THE UNKNOWN MORE CERTAIN: A STACKED ENSEMBLE CLASSIFIER FOR OPEN GESTURE RECOGNITION WITH A SOCIAL ROBOT	3918
<i>Heike Brock, Randy Gomez, Honda Research Institute Japan, Co Ltd., Japan</i>	
MLSP-29: SIGNAL PROCESSING THEORY AND METHODS	
MLSP-29.1: APPLYING DIFFERENTIAL PRIVACY TO TENSOR COMPLETION	3923
<i>Zheng Wei, Zhengpin Li, Jian Wang, Fudan University, China; Xiaojun Mao, Shanghai Jiao Tong University, China</i>	
MLSP-29.2: LOW-COMPLEXITY ATTENTION MODELLING VIA GRAPH TENSOR NETWORKS	3928
<i>Yao Lei Xu, Kriton Konstantinidis, Shengxi Li, Danilo P. Mandic, Imperial College London, United Kingdom of Great Britain and Northern Ireland; Ljubiša Stanković, University of Montenegro, Montenegro</i>	
MLSP-29.3: AN ACCELERATED RANK-(L,L,1,1) BLOCK TERM DECOMPOSITION OF MULTI-SUBJECT FMRI DATA UNDER SPATIAL ORTHONORMALITY CONSTRAINT	3933
<i>Li-Dan Kuang, Biao Wang, Hao-Peng Zhang, Jianming Zhang, Wenjun Li, Feng Li, Changsha University of Science and Technology, China; Qiu-Hua Lin, Dalian University of Technology, China; Vince D. Calhoun, Georgia State University, China</i>	
MLSP-29.5: IMPROVING DYNAMIC GRAPH CONVOLUTIONAL NETWORK WITH FINE-GRAINED ATTENTION MECHANISM	3938
<i>Bo Wu, Xun Liang, Xiangping Zheng, Yuhui Guo, Hui Tang, Renmin University of China, China</i>	
MLSP-29.6: ADAPID: AN ADAPTIVE PID OPTIMIZER FOR TRAINING DEEP NEURAL NETWORKS	3943
<i>Boxi Weng, Jian Sun, Gang Wang, School of Automation, Beijing Institute of Technology, China; Alireza Sadeghi, Dept. of Electrical and Computer Engineering, University of Minnesota, United States of America</i>	
MLSP-30: DEEP LEARNING III	
MLSP-30.1: MEMORY IN ECHO STATE NETWORKS AND THE CONTROLLABILITY MATRIX RANK	3948
<i>Brian Whiteaker, Peter Gerstoft, UC San Diego and Scripps Inst. of Oceanography, United States of America</i>	
MLSP-30.2: OT CLEANER: LABEL CORRECTION AS OPTIMAL TRANSPORT	3953
<i>Jun Xia, Zhejiang University, Westlake University, China; Cheng Tan, Lirong Wu, Yongjie Xu, Stan Z. Li, Westlake University, China</i>	

MLSP-30.3: DEMON: IMPROVED NEURAL NETWORK TRAINING WITH MOMENTUM DECAY	3958
<i>John Chen, Cameron Wolfe, Anastasios Kyrillidis, Rice University, United States of America; Zhao Li, UT Health, United States of America</i>	
MLSP-30.4: DEPTH PRUNING WITH AUXILIARY NETWORKS FOR TINYML	3963
<i>Josen Daniel De Leon, Rowel Atienza, University of the Philippines, Philippines</i>	
MLSP-30.5: GLASSOFORMER: A QUERY-SPARSE TRANSFORMER FOR POST-FAULT POWER GRID VOLTAGE PREDICTION	3968
<i>Yunling Zheng, Carson Hu, Jack Xin, University of California, Irvine, United States of America; Guang Lin, Purdue University, United States of America; Meng Yue, Brookhaven National Laboratory, United States of America; Bao Wang, University of Utah, United States of America</i>	
MLSP-30.6: SAR-SHIPNET: SAR-SHIP DETECTION NEURAL NETWORK VIA BIDIRECTIONAL COORDINATE ATTENTION AND MULTI-RESOLUTION FEATURE FUSION	3973
<i>Yuwen Deng, Donghai Guan, Weiwei Yuan, Jiemin Ji, Mingqiang Wei, Nanjing University of Aeronautics and Astronautics, China; Yanyu Chen, Fudan University, China</i>	
MLSP-31: MACHINE LEARNING FOR HEALTHCARE AND LIFE SCIENCES	
MLSP-31.1: SPATIO-TEMPORAL PRRS EPIDEMIC FORECASTING VIA FACTORIZED DEEP GENERATIVE MODELING	3978
<i>Mohammadsadegh Shamsabardeh, Beatriz Martínez-López, University of California, Davis, United States of America; Bahar Azari, Northeastern University, United States of America</i>	
MLSP-31.2: FUSION-ID: A PHOTOPLETHYSMOGRAPHY AND MOTION SENSOR FUSION BIOMETRIC AUTHENTICATOR WITH FEW-SHOT ON-BOARDING	3983
<i>Harshat Kumar, University of Pennsylvania, United States of America; Hojjat Seyed Mousavi, Behrooz Shahsavari, Apple Inc., United States of America</i>	
MLSP-31.3: DYNIMP: DYNAMIC IMPUTATION FOR WEARABLE SENSING DATA THROUGH SENSORY AND TEMPORAL RELATEDNESS	3988
<i>Zepeng Huo, Taowei Ji, Yifei Liang, Xiaoning Qian, Bobak Mortazavi, Texas A&M University, United States of America; Shuai Huang, University of Washington, United States of America; Zhangyang Wang, University of Texas at Austin, United States of America</i>	
MLSP-31.4: INCREMENTAL CONTEXT AWARE ATTENTIVE KNOWLEDGE TRACING	3993
<i>Cheryl Sze Yin Wong, Yang Guo, Nancy F. Chen, Savitha Ramasamy, Institute for Infocomm Research, Singapore</i>	
MLSP-31.5: ROBUST AND EFFICIENT UNCERTAINTY AWARE BIOSIGNAL CLASSIFICATION VIA EARLY EXIT ENSEMBLES	3998
<i>Alexander Campbell, Lorena Qendro, Pietro Liò, Cecilia Mascolo, University of Cambridge, United Kingdom of Great Britain and Northern Ireland</i>	
MLSP-31.6: TEMPORAL CROSS-GRAPH NETWORK FOR BRAIN FUNCTIONAL ACTIVITY PREDICTION	4003
<i>Xinyu Yuan, Wenhan Wang, Youyong Kong, Jiasong Wu, Guanyu Yang, Huazhong Shu, School of Computer Science and Engineering, Southeast University, China</i>	
MLSP-32: REINFORCEMENT LEARNING I	
MLSP-32.1: POPO: PESSIMISTIC OFFLINE POLICY OPTIMIZATION	4008
<i>Qiang He, Xinwen Hou, Yu Liu, Institute of Automation, Chinese Academy of Sciences, China</i>	
MLSP-32.2: BYZANTINE-ROBUST FEDERATED DEEP DETERMINISTIC POLICY GRADIENT	4013
<i>Qifeng Lin, Qing Ling, School of Computer Science and Engineering, Sun Yat-Sen University, China</i>	

MLSP-32.3: IMPROVING ACTOR-CRITIC REINFORCEMENT LEARNING VIA HAMILTONIAN MONTE CARLO METHOD	4018
<i>Duo Xu, Faramarz Fekri, Georgia Institute of Technology, United States of America</i>	
MLSP-32.4: EFFICIENT AND STABLE INFORMATION DIRECTED EXPLORATION FOR CONTINUOUS REINFORCEMENT LEARNING	4023
<i>Mingzhe Chen, Xi Xiao, Wanpeng Zhang, Tsinghua University, China; Xiaotian Gao, Microsoft Research Asia, China</i>	
MLSP-32.5: HYPERGRAPH-BASED REINFORCEMENT LEARNING FOR STOCK PORTFOLIO SELECTION	4028
<i>Xiaojie Li, Chaoran Cui, Donglin Cao, Juan Du, Chunyun Zhang, Shandong University of Finance and Economics, China</i>	
MLSP-33: DEEP LEARNING IV	
MLSP-33.1: MEMORY-BASED MESSAGE PASSING: DECOUPLING THE MESSAGE FOR PROPAGATION FROM DISCRIMINATION	4033
<i>Jie Chen, Weiqi Liu, Jian Pu, Fudan University, China</i>	
MLSP-33.2: PEAR: PHOTOGRAPHIC EMBEDDING FOR AESTHETIC RATING	4038
<i>Hao Wu, Jiangchao Yao, Alibaba Group, China; , ,</i>	
MLSP-33.3: GRADIENT-WEIGHTED CLASS ACTIVATION MAPPING FOR SPATIO TEMPORAL GRAPH CONVOLUTIONAL NETWORK	4043
<i>Pratyusha Das, Antonio Ortega, University of Southern California, United States of America</i>	
MLSP-33.4: DEEP LEARNING BASED OFF-ANGLE IRIS RECOGNITION	4048
<i>Ehsaneddin Jalilian, Georg Wimmer, Andreas Uhl, University of Salzburg, Austria; Mahmut Karakaya, Kennesaw State University, United States of America</i>	
MLSP-33.5: TOWARDS ROBUST VISUAL TRANSFORMER NETWORKS VIA K-SPARSE ATTENTION	4053
<i>Sajjad Amini, Shahrokh Ghaemmaghami, Sharif University of Technology, Iran (Islamic Republic of)</i>	
MLSP-33.6: A GLOBAL TO LOCAL GUIDING NETWORK FOR MISSING DATA IMPUTATION	4058
<i>Wei Wang, Yimeng Chai, Yue Li, nankai university, China</i>	
MLSP-34: MACHINE LEARNING FOR TELECOMMUNICATIONS	
MLSP-34.1: LOCUNET: FAST URBAN POSITIONING USING RADIO MAPS AND DEEP LEARNING	4063
<i>Çağkan Yapar, Giuseppe Caire, Technische Universität Berlin, Germany; Ron Levie, Gitta Kutyniok, Ludwig-Maximilians-Universität München, Germany</i>	
MLSP-34.2: LITEHAR: LIGHTWEIGHT HUMAN ACTIVITY RECOGNITION FROM WIFI SIGNALS WITH RANDOM CONVOLUTION KERNELS	4068
<i>Hojjat Salehinejad, Shahrokh Valaee, University of Toronto, Canada</i>	
MLSP-34.3: CDX-NET: CROSS-DOMAIN MULTI-FEATURE FUSION MODELING VIA DEEP NEURAL NETWORKS FOR MULTIVARIATE TIME SERIES FORECASTING IN AIOPS	4073
<i>Jiajia Li, Ling Dai, Bin Sheng, Shanghai Jiao Tong University, China; Feng Tan, Zikai Wang, Shanghai Artificial Intelligence Research Institute, China; Hui Shen, Shanghai Dingmao information technology Inc., China; Pengwei Hu, Merck China Innovation Hub, China</i>	
MLSP-34.4: A CLUSTERING-BASED ML SCHEME FOR CAPACITY APPROACHING SOFT LEVEL SENSING IN 3D TLC NAND	4078
<i>Li-Wei Liu, Yen-Ching Liao, Hsie-Chia Chang, National Yang Ming Chiao Tung University, Taiwan</i>	

MLSP-34.5: DYNAMIC RESOURCE OPTIMIZATION FOR ADAPTIVE FEDERATED LEARNING EMPOWERED BY RECONFIGURABLE INTELLIGENT SURFACES	4083
<i>Claudio Battiloro, Paolo Di Lorenzo, Sergio Barbarossa, Sapienza University of Rome, Italy; Mattia Merluzzi, CEA-Leti, Université Grenoble Alpes, France</i>	
MLSP-34.6: LEARNING-BASED RESOURCE ALLOCATION WITH DYNAMIC DATA RATE CONSTRAINTS	4088
<i>Pourya Behmandpoor, Panagiotis Patrinos, Marc Moonen, KU Leuven, Belgium</i>	
MLSP-35: REINFORCEMENT LEARNING II	
MLSP-35.1: DENOISING-ORIENTED DEEP HIERARCHICAL REINFORCEMENT LEARNING FOR NEXT-BASKET RECOMMENDATION	4093
<i>Qihan Du, Li Yu, Huiyuan Li, Youfang Leng, Ningrui Ou, Renmin University of China, China</i>	
MLSP-35.2: COMPETITIVE MULTI-AGENT REINFORCEMENT LEARNING WITH SELF-SUPERVISED REPRESENTATION	4098
<i>DiJia Su, Jason D. Lee, John M. Mulvey, H. Vincent Poor, Princeton University, United States of America</i>	
MLSP-35.3: MODEL-BASED ONLINE LEARNING FOR RESOURCE SHARING IN JOINT RADAR-COMMUNICATION SYSTEMS	4103
<i>Petteri Pulkkinen, Aalto University, Saab Finland Oy, Finland; Visa Koivunen, Aalto University, Finland</i>	
MLSP-35.4: QRELATION: AN AGENT RELATION-BASED APPROACH FOR MULTI-AGENT REINFORCEMENT LEARNING VALUE FUNCTION FACTORIZATION	4108
<i>Siqi Shen, Jun Liu, Mengwei Qiu, Weiquan Liu, Cheng Wang, Xiamen University, China; Yongquan Fu, Qinglin Wang, Peng Qiao, National University of Defense Technology, China, China</i>	
MLSP-35.5: DENOISING-GUIDED DEEP REINFORCEMENT LEARNING FOR SOCIAL RECOMMENDATION	4113
<i>Qihan Du, Li Yu, Huiyuan Li, Youfang Leng, Ningrui Ou, Junyao Xiang, Renmin University of China, China</i>	
MLSP-36: APPLICATIONS OF MACHINE LEARNING	
MLSP-36.1: AN EFFICIENT DP-SGD MECHANISM FOR LARGE SCALE NLU MODELS	4118
<i>Christophe Dupuy, Radhika Arava, Rahul Gupta, Amazon Alexa AI, United States of America; Anna Rumshisky, Amazon Alexa AI; University of Massachusetts, Lowell, United States of America</i>	
MLSP-36.2: MAKD:MULTIPLE AUXILIARY KNOWLEDGE DISTILLATION	4123
<i>Zehan Chen, Xuan Jin, Yuan He, Hui Xue, Alibaba Group, China</i>	
MLSP-36.3: FEATURE IMITATING NETWORKS	4128
<i>Sari Saba-Sadiya, Mohammad Ghassemi, Michigan State University, United States of America; Tuka Alhanai, New York University Abu Dhabi, United Arab Emirates</i>	
MLSP-36.5: OVER-PARAMETERIZED NETWORK SOLVES PHASE RETRIEVAL EFFECTIVELY	4133
<i>Ji Li, Chao Wang, National University of Singapore, Singapore</i>	
MLSP-36.6: DEEP SPATIO-TEMPORAL WIND POWER FORECASTING	4138
<i>Jiangyuan Li, Mohammadreza Armandpour, Texas A&M University, United States of America</i>	
MLSP-37: KERNEL METHODS	
MLSP-37.1: MULTIPLE KERNEL K-MEANS CLUSTERING WITH SIMULTANEOUS SPECTRAL ROTATION	4143
<i>Jitao Lu, Yihang Lu, Rong Wang, Feiping Nie, Xuelong Li, Northwestern Polytechnical University, China</i>	

MLSP-37.2: MULTITASK GAUSSIAN PROCESS WITH HIERARCHICAL LATENT INTERACTIONS	4148
<i>Kai Chen, Feng Yin, Shuguang Cui, The Chinese University of Hong Kong, Shenzhen, China; Twan van Laarhoven, Elena Marchiori, Radboud University, Nijmegen, Netherlands</i>	
MLSP-37.3: DISCRETE MULTI-KERNEL K-MEANS WITH DIVERSE AND OPTIMAL KERNEL LEARNING	4153
<i>Yihang Lu, Jitao Lu, Rong Wang, Feiping Nie, Northwestern Polytechnical University, China</i>	
MLSP-37.4: ACCESS CONTROL FOR PRIVACY-PRESERVING GAUSSIAN PROCESS REGRESSION	4158
<i>Takayuki Nakachi, University of the Ryukyus, Japan; Yitu Wang, Nippon Telegraph and Telephone Corporation, Japan</i>	
MLSP-37.5: SCALABLE RIDGE LEVERAGE SCORE SAMPLING FOR THE NYSTRÖM METHOD	4163
<i>Farah Cherfaoui, Hachem Kadri, LIS, France; Liva Ralaivola, Criteo IA lab, Paris, France, France</i>	
MLSP-37.6: ON SUBMODULAR SET COVER PROBLEMS FOR NEAR-OPTIMAL ONLINE KERNEL BASIS SELECTION	4168
<i>Hrusikesh Pradhan, IIT Kanpur, India; Alec Koppel, Amazon, United States of America; Ketan Rajawat, Indian Institute of Technology Kanpur, India</i>	
MLSP-38: SEQUENTIAL LEARNING	
MLSP-38.1: IMPROVING FEATURE GENERALIZABILITY WITH MULTITASK LEARNING IN CLASS INCREMENTAL LEARNING	4173
<i>Dong Ma, Chi Ian Tang, Cecilia Mascolo, University of Cambridge, United Kingdom of Great Britain and Northern Ireland</i>	
MLSP-38.2: ADVERSARIAL MASK TRANSFORMER FOR SEQUENTIAL LEARNING	4178
<i>Hou Lio, Shang-En Li, Jen-Tzung Chien, National Yang Ming Chiao Tung University, Taiwan</i>	
MLSP-38.3: ONLINE LEARNING WITH PROBABILISTIC FEEDBACK	4183
<i>Pouya M. Ghari, Yanning Shen, University of California Irvine, United States of America</i>	
MLSP-38.4: DATA INCUBATION — SYNTHESIZING MISSING DATA FOR HANDWRITING RECOGNITION	4188
<i>Jen-Hao Rick Chang, Martin Bresler, Youssouf Chherawala, Adrien Delaye, Thomas Deselaers, Ryan Dixon, Oncel Tuzel, Apple, United States of America</i>	
MLSP-38.5: TRACKING THE DIMENSIONS OF LATENT SPACES OF GAUSSIAN PROCESS LATENT VARIABLE MODELS	4193
<i>Yuhao Liu, Petar M. Djuric, Stony Brook University, United States of America</i>	
MLSP-38.6: CONTROLLED SENSING AND ANOMALY DETECTION VIA SOFT ACTOR-CRITIC REINFORCEMENT LEARNING	4198
<i>Chen Zhong, M. Cenk Gursoy, Senem Velipasalar, Syracuse University, United States of America</i>	
MLSP-39: DATA SCIENCE INITIATIVE II	
MLSP-40: DEEP LEARNING V	
MLSP-40.1: WIN THE LOTTERY TICKET VIA FOURIER ANALYSIS: FREQUENCIES GUIDED NETWORK PRUNING	4203
<i>Yuzhang Shang, Bin Duan, Yan Yan, Illinois Institute of Technology, United States of America; Ziliang Zong, Texas State University, United States of America; Liqiang Nie, Shandong University, China; , ,</i>	

MLSP-40.2: SPARSEBFA: ATTACKING SPARSE DEEP NEURAL NETWORKS WITH THE WORST-CASE BIT FLIPS ON COORDINATES	4208
<i>Kyungmi Lee, Anantha P. Chandrakasan, Massachusetts Institute of Technology, United States of America</i>	
MLSP-40.4: ADVERSARIAL EXAMPLES DETECTION BASED ON ERROR LEVEL ANALYSIS AND SPACE MAPPING	4213
<i>Sizhao Huang, Shuai Wang, Jian Chen, Guozhi Li, School of Information and Communication Engineering University of Electronic Science and Technology of China; Yangtze Delta Region Institute of University of Electronic Science and Technology of China, China; Wenyi Wang, School of Information and Communication Engineering University of Electronic Science and Technology of China, China</i>	
MLSP-40.5: LEARNING MONOCULAR 3D HUMAN POSE ESTIMATION WITH SKELETAL INTERPOLATION	4218
<i>Ziyi Chen, Shang-Hong Lai, National Tsing Hua University, Taiwan; Akihiro Sugimoto, National Institute of Informatics, Japan</i>	
MLSP-40.6: TRAINING STABLE GRAPH NEURAL NETWORKS THROUGH CONSTRAINED LEARNING	4223
<i>Juan Cervino, Luana Ruiz, Alejandro Ribeiro, University of Pennsylvania, United States of America</i>	
MLSP-41: SUPERVISED LEARNING	
MLSP-41.1: MISMATCHED SUPERVISED LEARNING	4228
<i>Xun Xian, Mingyi Hong, Jie Ding, University of Minnesota, United States of America</i>	
MLSP-41.2: SUPERVISED TRAINING OF SIAMESE SPIKING NEURAL NETWORKS WITH EARTH MOVER'S DISTANCE	4233
<i>Mateusz Pabian, Dominik Rzepka, Mirosław Pawlak, AGH University of Science and Technology, Poland</i>	
MLSP-41.3: ON THE EFFECTIVENESS OF ACTIVE LEARNING BY UNCERTAINTY SAMPLING IN CLASSIFICATION OF HIGH-DIMENSIONAL GAUSSIAN MIXTURE DATA	4238
<i>Xiaoyi Mai, French National Centre for Scientific Research (CNRS), France; Salman Avestimehr, Antonio Ortega, Mahdi Soltanolkotabi, University of Southern California, United States of America</i>	
MLSP-41.4: NEURAL COLLAPSE IN DEEP HOMOGENEOUS CLASSIFIERS AND THE ROLE OF WEIGHT DECAY	4243
<i>Akshay Rangamani, Andrzej Banburski-Fahey, Massachusetts Institute of Technology, United States of America</i>	
MLSP-41.5: SYNTHESIS OF ADVERSARIAL SAMPLES IN TWO-STAGE CLASSIFIERS	4248
<i>Ismail Alkhouri, George Atia, University of Central Florida, United States of America; Alvaro Velasquez, Air Force Research Laboratory, United States of America</i>	
MLSP-41.6: SYNERGISTIC NETWORK LEARNING AND LABEL CORRECTION FOR NOISE-ROBUST IMAGE CLASSIFICATION	4253
<i>Chen Gong, University of Washington, United States of America; Kong Bin, Xin Wang, Youbing Yin, Qi Song, Keya medical, United States of America; Eric Seibel, University of Washington, United States of America</i>	
MLSP-42: DISTRIBUTED AND FEDERATED LEARNING I	
MLSP-42.1: SOCIAL WELFARE MAXIMIZATION IN CROSS-SILO FEDERATED LEARNING	4258
<i>Jianan Chen, Qin Hu, Indiana University Purdue University Indianapolis, United States of America; Honglu Jiang, The University of Texas Rio Grande Valley, United States of America</i>	
MLSP-42.2: PRIVACY-PRESERVING DISTRIBUTED EXPECTATION MAXIMIZATION FOR GAUSSIAN MIXTURE MODEL USING SUBSPACE PERTURBATION	4263
<i>Qiongxiu Li, Jaron Skovsted Gundersen, Katrine Tjell, Rafal Wisniewski, Mads Græsbøll Christensen, Aalborg university, Denmark</i>	

MLSP-42.3: A COMMUNICATION EFFICIENT QUASI-NEWTON METHOD FOR LARGE-SCALE DISTRIBUTED MULTI-AGENT OPTIMIZATION	4268
<i>Yichuan Li, University of Illinois Urbana Champaign, United States of America; Petros Voulgaris, University of Nevada, Reno, United States of America; Nikolaos Freris, University of Science and Technology of China, China</i>	
MLSP-42.4: A BYZANTINE-RESILIENT DUAL SUBGRADIENT METHOD FOR VERTICAL FEDERATED LEARNING	4273
<i>Kun Yuan, Alibaba Group (US), United States of America; Zhaoxian Wu, Qing Ling, Sun Yat-Sen University, China</i>	
MLSP-42.5: BYZANTINE-ROBUST AGGREGATION WITH GRADIENT DIFFERENCE COMPRESSION AND STOCHASTIC VARIANCE REDUCTION FOR FEDERATED LEARNING	4278
<i>Heng Zhu, University of California, San Diego, United States of America; Qing Ling, Sun Yat-Sen University, China</i>	
MLSP-42.6: VARIANCE REDUCTION-BOOSTED BYZANTINE ROBUSTNESS IN DECENTRALIZED STOCHASTIC OPTIMIZATION	4283
<i>Jie Peng, Qing Ling, Sun Yat-Sen University, China; Weiyu Li, Harvard Univeristy, China</i>	
MLSP-43: DEEP LEARNING FOR SPEECH AND LANGUAGE PROCESSING	
MLSP-43.1: INTEGER-ONLY ZERO-SHOT QUANTIZATION FOR EFFICIENT SPEECH RECOGNITION	4288
<i>Sehoon Kim, Amir Gholami, Zhewei Yao, Nicholas Lee, Patrick Wang, Aniruddha Nrusimha, Bohan Zhai, Tianren Gao, Michael Mahoney, Kurt Keutzer, University of California, Berkeley, United States of America</i>	
MLSP-43.3: NNSPEECH: SPEAKER-GUIDED CONDITIONAL VARIATIONAL AUTOENCODER FOR ZERO-SHOT MULTI-SPEAKER TEXT-TO-SPEECH	4293
<i>Botao Zhao, Fudan University, China; Xulong Zhang, Jianzong Wang, Ning Cheng, Jing Xiao, Ping An Technology (Shenzhen) Co., Ltd., China</i>	
MLSP-43.4: NOISE-ROBUST SPEECH RECOGNITION WITH 10 MINUTES UNPARALLELED IN-DOMAIN DATA	4298
<i>Chen Chen, Nana Hou, Yuchen Hu, Eng Siong Chng, Nanyang Technological University, Singapore; Shashank Shirol, Manipal Institute of Technology, India</i>	
MLSP-43.5: ENHANCING CLASS UNDERSTANDING VIA PROMPT-TUNING FOR ZERO-SHOT TEXT CLASSIFICATION	4303
<i>Yuhao Dan, Qin Chen, Liang He, East China Normal University, China; Jie Zhou, Fudan Univeristy, China; Qingchun Bai, Shanghai Open University, China</i>	
MLSP-43.6: FILTERAUGMENT: AN ACOUSTIC ENVIRONMENTAL DATA AUGMENTATION METHOD	4308
<i>Hyeonuk Nam, Seong-Hu Kim, Yong-Hwa Park, KAIST, Korea, Republic of</i>	
MLSP-44: LEARNING THEORY AND ALGORITHMS I	
MLSP-44.1: THE REPRESENTATION JENSEN-RÉNYI DIVERGENCE	4313
<i>Jhoan Keider Hoyos Osorio, Oscar Skean, Luis Gonzalo Sanchez Giraldo, University of Kentucky, United States of America; Austin J. Brockmeier, University of Delaware, United States of America</i>	
MLSP-44.3: MULTI-VIEW INFORMATION BOTTLENECK WITHOUT VARIATIONAL APPROXIMATION	4318
<i>Qi Zhang, Jingmin Xin, Badong Chen, Xi'an Jiaotong University, China; Shujian Yu, UiT - The Arctic University of Norway, Norway</i>	
MLSP-44.4: TIME-FREQUENCY AND GEOMETRIC ANALYSIS OF TASK-DEPENDENT LEARNING IN RAW WAVEFORM BASED ACOUSTIC MODELS	4323
<i>Devansh Gupta, Vinayak Abrol, Indraprastha Institute of Information Technology, Delhi, India</i>	

MLSP-44.6: CHANNEL REDUNDANCY AND OVERLAP IN CONVOLUTIONAL NEURAL NETWORKS WITH CHANNEL-WISE NNK GRAPHS	4328
<i>David Bonet, Javier Ruiz-Hidalgo, Universitat Politècnica de Catalunya, Spain; Antonio Ortega, Sarath Shekkizhar, University of Southern California, United States of America</i>	
 MLSP-45: DISTRIBUTED AND FEDERATED LEARNING II	
MLSP-45.1: FEDCLEAN: A DEFENSE MECHANISM AGAINST PARAMETER POISONING ATTACKS IN FEDERATED LEARNING	4333
<i>Abhishek Kumar, University of Oulu & University of Helsinki, Finland; Vivek Khimani, Drexel University, United States of America; Dimitris Chatzopoulos, University College Dublin, Ireland; Pan Hui, Hong Kong University of Science and Technology & University of Helsinki, Hong Kong</i>	
MLSP-45.2: A METHOD TO REVEAL SPEAKER IDENTITY IN DISTRIBUTED ASR TRAINING, AND HOW TO COUNTER IT	4338
<i>Trung Dang, Peter Chin, Boston University, United States of America; Om Thakkar, Swaroop Ramaswamy, Rajiv Mathews, Françoise Beaufays, Google LLC, United States of America</i>	
MLSP-45.3: ON THE CONVERGENCE OF ADAM-TYPE ALGORITHMS FOR SOLVING STRUCTURED SINGLE NODE AND DECENTRALIZED MIN-MAX SADDLE POINT GAMES	4343
<i>Babak Barazandeh, Kristal Curtis, Chandrima Sarkar, Splunk, United States of America; Ram Sriharsha, Pinecone, United States of America; George Michailidis, University of Florida, United States of America</i>	
MLSP-45.4: PARTIAL VARIABLE TRAINING FOR EFFICIENT ON-DEVICE FEDERATED LEARNING	4348
<i>Tien-Ju Yang, Dhruv Guliani, Françoise Beaufays, Giovanni Motta, Google LLC, United States of America</i>	
MLSP-45.5: GRADIENT STALENESS IN ASYNCHRONOUS OPTIMIZATION UNDER RANDOM COMMUNICATION DELAYS	4353
<i>Haider Al-Lawati, Stark Draper, University of Toronto, Canada</i>	
MLSP-45.6: TEMPO: IMPROVING TRAINING PERFORMANCE IN CROSS-SILO FEDERATED LEARNING	4358
<i>Chen Ying, Baochun Li, University of Toronto, Canada; Bo Li, Hong Kong University of Science and Technology, Hong Kong</i>	
 MLSP-46: SOURCE SEPARATION AND ICA	
MLSP-46.1: DMANET: DEEP LEARNING-BASED DIFFERENTIAL MICROPHONE ARRAYS FOR MULTI-CHANNEL SPEECH SEPARATION	4363
<i>Xiaokang Yang, Jianguo Wei, Tianjin University, China</i>	
MLSP-46.2: AMICABLE EXAMPLES FOR INFORMED SOURCE SEPARATION.....	4368
<i>Naoya Takahashi, Yuki Mitsufuji, Sony Group Corporation, Japan</i>	
MLSP-46.3: REMIX-CYCLE-CONSISTENT LEARNING ON ADVERSARIALY LEARNED SEPARATOR FOR ACCURATE AND STABLE UNSUPERVISED SPEECH SEPARATION	4373
<i>Kohei Saijo, Tetsuji Ogawa, Waseda University, Japan</i>	
MLSP-46.4: AN INFORMATION MAXIMIZATION BASED BLIND SOURCE SEPARATION APPROACH FOR DEPENDENT AND INDEPENDENT SOURCES	4378
<i>Alper Erdogan, Koc University, Turkey</i>	
MLSP-46.5: BLIND SEPARATION OF LINEAR-QUADRATIC MIXTURES OF MUTUALLY INDEPENDENT AND AUTOCORRELATED SOURCES	4383
<i>Shahram Hosseini, Yannick Deville, Toulouse university, France</i>	

MLSP-46.6: LARGE-SCALE INDEPENDENT COMPONENT ANALYSIS BY SPEEDING UP LIE GROUP TECHNIQUES	4388
<i>Matthias Hermann, Georg Umlauf, Matthias O. Franz, HTWG Konstanz, Germany</i>	
MLSP-47: DEEP LEARNING VI	
MLSP-47.1: PREDICTING THE GENERALIZATION GAP IN DEEP MODELS USING ANCHORING	4393
<i>Vivek Narayanaswamy, Andreas Spanias, Arizona State University, United States of America; Rushil Anirudh, Irene Kim, Yamen Mubarka, Jayaraman J Thiagarajan, Lawrence Livermore National Laboratory, United States of America</i>	
MLSP-47.2: WHEN DOES BACKDOOR ATTACK SUCCEED IN IMAGE RECONSTRUCTION? A STUDY OF HEURISTICS VS. BI-LEVEL SOLUTION	4398
<i>Vardaan Taneja, Indian Institute of Technology, Delhi, India; Pin-Yu Chen, IBM Research, United States of America; Yuguang Yao, Sijia Liu, Michigan State University, United States of America</i>	
MLSP-47.3: MIXED KNOWLEDGE RELATION TRANSFORMER FOR IMAGE CAPTIONING	4403
<i>Tianyu Chen, Zhixin Li, Jiahui Wei, Tiantao Xian, Guangxi Normal University, China</i>	
MLSP-47.4: BALANCED STRIPE-WISE PRUNING IN THE FILTER	4408
<i>Zheng Huo, Chong Wang, Weiwei Chen, Yuqi Li, Ningbo University, China; Jun Wang, China University of Mining and Technology, China; Jiafei Wu, SenseTime Research, China</i>	
MLSP-47.5: GAN-BASED JOINT ACTIVITY DETECTION AND CHANNEL ESTIMATION FOR GRANT-FREE RANDOM ACCESS	4413
<i>Shuang Liang, Yinan Zou, ShanghaiTech University, Shanghai Institute of Microsystem and Information Technology and University of Chinese Academy of Sciences, China; Yong Zhou, ShanghaiTech University, China</i>	
MLSP-48: LEARNING THEORY AND ALGORITHMS II	
MLSP-48.1: CASCADING BANDIT UNDER DIFFERENTIAL PRIVACY	4418
<i>Kun Wang, Shuai Li, Shanghai Jiao Tong University, China; Jing Dong, University of Michigan, China; Baoxiang Wang, The Chinese University of Hong Kong, Shenzhen, China</i>	
MLSP-48.2: ITERATIVE RE-WEIGHTED LEAST SQUARES ALGORITHMS FOR NON-NEGATIVE SPARSE AND GROUP-SPARSE RECOVERY	4423
<i>Angshul Majumdar, Indraprastha Institute of Information Technology, India</i>	
MLSP-48.3: EXACT PARTITIONING OF HIGH-ORDER PLANTED MODELS WITH A TENSOR NUCLEAR NORM CONSTRAINT	4428
<i>Chuyang Ke, Jean Honorio, Purdue University, United States of America</i>	
MLSP-48.4: NO MORE THAN 6FT APART: ROBUST K-MEANS VIA RADIUS UPPER BOUNDS	4433
<i>Ahmed Imtiaz Humayun, Anastasios Kyrillidis, Richard Baraniuk, Rice University, United States of America; Randall Balestriero, Meta AI Research, United States of America</i>	
MLSP-48.5: DEEP KERNEL LEARNING NETWORKS WITH MULTIPLE LEARNING PATHS	4438
<i>Ping Xu, Yue Wang, Xiang Chen, Zhi Tian, George Mason University, United States of America</i>	
MLSP-48.6: PROVABLE SAMPLE COMPLEXITY GUARANTEES FOR LEARNING OF CONTINUOUS-ACTION GRAPHICAL GAMES WITH NONPARAMETRIC UTILITIES	4443
<i>Adarsh Barik, Jean Honorio, Purdue University, United States of America</i>	

MLSP-49: LEARNING FROM MULTIMODAL DATA

MLSP-49.1: CROSS-MODAL KNOWLEDGE DISTILLATION FOR VISION-TO-SENSOR ACTION RECOGNITION 4448

Jianyuan Ni, Anne H. H. Ngu, Texas State University, United States of America; Raunak Sarbajna, University of Houston, United States of America; Yang Liu, Sun Yat-sen University, China; Yan Yan, Illinois Institute of Technology, United States of America

MLSP-49.2: CLIPCAM: A SIMPLE BASELINE FOR ZERO-SHOT TEXT-GUIDED OBJECT AND ACTION LOCALIZATION 4453

Hsuan-An Hsia, Che-Hsien Lin, Bo-Han Kung, Jhao-Ting Chen, Daniel Stanley Tan, Jun-Cheng Chen, Academia Sinica, Taiwan; Kai-Lung Hua, National Taiwan University of Science and Technology, Taiwan

MLSP-49.3: EXPLORING DUAL STREAM GLOBAL INFORMATION FOR IMAGE CAPTIONING 4458

Tiantao Xian, Zhixin Li, Tianyu Chen, Guangxi Normal University, China; Huifang Ma, Northwest Normal University, China

MLSP-49.4: UNSUPERVISED CONTRASTIVE HASHING FOR CROSS-MODAL RETRIEVAL IN REMOTE SENSING 4463

Georgii Mikriukov, Mahdyar Ravanbakhsh, Begüm Demir, Technische Universität Berlin, Germany

MLSP-49.5: ROBUST THERMAL INFRARED PEDESTRIAN DETECTION BY ASSOCIATING VISIBLE PEDESTRIAN KNOWLEDGE 4468

Sungjune Park, Dae Hwi Choi, Jung Uk Kim, Yong Man Ro, Korea Advanced Institute of Science and Technology (KAIST), Korea, Republic of

MLSP-49.6: A GENERALIZED HIERARCHICAL NONNEGATIVE TENSOR DECOMPOSITION 4473

Joshua Vendrow, Deanna Needell, UCLA, United States of America; Jamie Haddock, Harvey Mudd College, United States of America

MLSP-50: DEEP LEARNING IV

MLSP-50.1: TWO STRATEGIES TOWARD LIGHTWEIGHT IMAGE SUPER-RESOLUTION 4478

Zongcai Du, Jie Liu, Jie Tang, Gangshan Wu, Nanjing University, China

MLSP-50.2: ON MINI-BATCH TRAINING WITH VARYING LENGTH TIME SERIES..... 4483

Brian Kenji Iwana, Kyushu University, Japan

MLSP-50.3: ACP: ADAPTIVE CHANNEL PRUNING FOR EFFICIENT NEURAL NETWORKS 4488

Yuan Zhang, Yuan Yuan, Qi Wang, Northwestern Polytechnical University, China

MLSP-50.4: BAYESIAN CONTINUAL IMPUTATION AND PREDICTION FOR IRREGULARLY SAMPLED TIME SERIES DATA 4493

Yang Guo, Cheryl Sze Yin Wong, Savitha Ramasamy, Institute for Infocomm Research, A*STAR, Singapore; Jeanette Wen Jun Poh, Nanyang Technological University Singapore, Singapore

MLSP-50.5: CONFIDENCE-AWARE MULTI-TEACHER KNOWLEDGE DISTILLATION..... 4498

Hailin Zhang, Defang Chen, Can Wang, Zhejiang University, China

MLSP-50.6: LEARNABLE HYPERGRAPH LAPLACIAN FOR HYPERGRAPH LEARNING 4503

Jiying Zhang, Yuzhao Chen, Xi Xiao, Runiu Lu, Shu-Tao Xia, Tsinghua University, China

MLSP-51: LEARNING THEORY AND ALGORITHMS III

MLSP-51.1: GRAPH LEARNING FROM MULTIVARIATE DEPENDENT TIME SERIES 4508 VIA A MULTI-ATTRIBUTE FORMULATION

Jitendra Tugnait, Auburn University, United States of America

MLSP-51.2: GENERALIZED SLICED PROBABILITY METRICS 4513

Soheil Kolouri, Vanderbilt University, United States of America; Kimia Nadjahi, Telecom Paris, France; Shahin Shahrampour, Northeastern University, United States of America; Umut Simsekli, INRIA / ENS, France

MLSP-51.3: RECOVERY OF NOISY POOLED TESTS VIA LEARNED FACTOR GRAPHS 4518 WITH APPLICATION TO COVID-19 TESTING

Eyal Fishel Ben-Knaan, Nir Shlezinger, Ben Gurion University, Israel; Yonina Eldar, Weizmann Institute of Science, Israel

MLSP-51.4: A REMEDY FOR DISTRIBUTIONAL SHIFTS THROUGH EXPECTED 4523 DOMAIN TRANSLATION

Jean-Christophe Gagnon-Audet, Irina Rish, University of Montréal and Mila - Québec AI Institute, Canada; Soroosh Shahtalebi, Vector Institute for Artificial intelligence, Canada; Frank Rudzicz, University of Toronto and Vector Institute for Artificial intelligence, Canada

MLSP-51.5: DETERMINISTIC TRANSFORM BASED WEIGHT MATRICES FOR 4528 NEURAL NETWORKS

Pol Grau Jurado, Saikat Chatterjee, KTH Royal Institute of Technology, Sweden; Xinyue Liang, Ocean University of China, China

MLSP-51.6: ADAPTIVE GROUP TESTING WITH MISMATCHED MODELS 4533

Mingzhou Fan, Byung-Jun Yoon, Edward R. Dougherty, Xiaoning Qian, Texas A&M University, United States of America; Francis Alexander, Brookhaven National Laboratory, United States of America

MLSP-52: MULTIMODAL DATA FUSION AND PROCESSING

MLSP-52.1: MIXED IN TIME AND MODALITY: CURSE OR BLESSING? 4538 CROSS-INSTANCE DATA AUGMENTATION FOR WEAKLY SUPERVISED MULTIMODAL TEMPORAL FUSION

Yonggang Zhu, Beijing University of Posts and Telecommunications, Peking University, China; Chao Tian, China Xi'an Satellite Control Center, China; Zhuqing Jiang, Aidong Men, Haiying Wang, Beijing University of Posts and Telecommunications, China; Qingchao Chen, Peking University, China

MLSP-52.2: MTAf: SHOPPING GUIDE MICRO-VIDEOS POPULARITY PREDICTION 4543 USING MULTIMODAL AND TEMPORAL ATTENTION FUSION APPROACH

Ningrui Ou, Li Yu, Huiyuan Li, Qihan Du, Junyao Xiang, Wei Gong, Renmin University of China, China

MLSP-52.3: LEARNING TO INTEGRATE VISION DATA INTO ROAD NETWORK DATA 4548

Oliver Stromann, Linköping University, Scania CV AB, Sweden; Alireza Razavi, Scania CV AB, Sweden; Michael Felsberg, Linköping University, Sweden

MLSP-52.4: HIERARCHICAL SIGNAL FUSION NETWORK FOR PULSAR DETECTION 4553 WITH PHASE-CORRELATION AND SIGNAL ATTENTIONS

Huajian Wu, Mingmin Chi, Fudan University, China

MLSP-52.5: RECOGNITION OF SILENTLY SPOKEN WORD FROM EEG SIGNALS 4558 USING DENSE ATTENTION NETWORK (DAN).

Sahil Datta, Akuha Aondoakaa, Jorunn Jo Holmberg, Elena Antonova, Brunel University London, United Kingdom of Great Britain and Northern Ireland

MLSP-53: MULTIMODAL ANALYSIS IN AUDIO APPLICATIONS

MLSP-53.1: WAV2CLIP: LEARNING ROBUST AUDIO REPRESENTATIONS FROM 4563 CLIP

Ho-Hsiang Wu, Juan Bello, New York University, United States of America; Prem Seetharaman, Kundan Kumar, Descript, United States of America

MLSP-53.2: ASD-TRANSFORMER: EFFICIENT ACTIVE SPEAKER DETECTION USING 4568 SELF AND MULTIMODAL TRANSFORMERS

Gourav Datta, University of Southern California, United States of America; Tyler Etchart, Vivek Yadav, Varsha Hedau, Pradeep Natarajan, Shih-Fu Chang, Amazon, United States of America

MLSP-53.3: MMLATCH: BOTTOM-UP TOP-DOWN FUSION FOR MULTIMODAL 4573 SENTIMENT ANALYSIS

Georgios Paraskevopoulos, Efthymios Georgiou, Alexandros Potamianos, National Technical University of Athens, Greece

MLSP-53.4: MULTI-CHANNEL ATTENTIVE GRAPH CONVOLUTIONAL NETWORK 4578 WITH SENTIMENT FUSION FOR MULTIMODAL SENTIMENT ANALYSIS

Luwei Xiao, Xingjiao Wu, Wen Wu, Jing Yang, Liang He, East China Normal University, China

MLSP-53.5: LEARNING MUSIC SEQUENCE REPRESENTATION FROM TEXT 4583 SUPERVISION

Tianyu Chen, Shuai Zhang, Haoyi Zhou, Jianxin Li, Beihang University, China; Yuan Xie, The Institute of Acoustics of the Chinese Academy of Sciences, China; Shaohan Huang, Microsoft, China

MLSP-53.6: ENHANCING AFFECTIVE REPRESENTATIONS OF MUSIC-INDUCED 4588 EEG THROUGH MULTIMODAL SUPERVISION AND LATENT DOMAIN ADAPTATION

Kleanthis Avramidis, University of Southern California, United States of America; Christos Garoufis, Athanasia Zlatintsi, Petros Maragos, National Technical University of Athens, Greece

MLSP-54: NEURAL AUDIO REPRESENTATION AND SYNTHESIS

MLSP-54.1: TOWARDS LEARNING UNIVERSAL AUDIO REPRESENTATIONS 4593

Luyu Wang, Pauline Luc, Yan Wu, Adria Recasens, Lucas Smaira, Andrew Brock, Andrew Jaegle, Jean-Baptiste Alayrac, Sander Dieleman, Joao Carreira, Aaron van den Oord, DeepMind, United Kingdom of Great Britain and Northern Ireland

MLSP-54.2: DIFFERENTIABLE WAVETABLE SYNTHESIS 4598

Siyuan Shan, University of North Carolina at Chapel Hill, United States of America; Lamtharn Hantrakul, Jitong Chen, Matt Avent, David Trevelyan, ByteDance, Thailand

MLSP-54.3: NEURAL AUDIO-TO-SCORE MUSIC TRANSCRIPTION FOR 4603 UNCONSTRAINED POLYPHONY USING COMPACT OUTPUT REPRESENTATIONS

Víctor Arroyo, Jose J. Valero-Mas, Jorge Calvo-Zaragoza, Antonio Pertusa, University of Alicante, Spain

MLSP-54.4: END-TO-END MUSIC REMASTERING SYSTEM USING 4608 SELF-SUPERVISED AND ADVERSARIAL TRAINING

Junghyun Koo, Seungryeol Paik, Kyogu Lee, Seoul National University, Korea, Republic of

MLSP-54.5: AVQVC: ONE-SHOT VOICE CONVERSION BY VECTOR 4613 QUANTIZATION WITH APPLYING CONTRASTIVE LEARNING

Huaizhen Tang, University of Science and Technology of China, China; Xulong Zhang, Jianzong Wang, Ning Cheng, Jing Xiao, Ping An Technology (Shenzhen) Co., Ltd., China

MLSP-54.6: TOWARDS SPEAKER AGE ESTIMATION WITH LABEL DISTRIBUTION 4618 LEARNING

Shijing Si, Jianzong Wang, Junqing Peng, Jing Xiao, Ping An Technology (Shenzhen) Co., Ltd., China

MMSP-1: MULTIMODAL SIGNAL PROCESSING, ANALYSIS, AND SYNTHESIS I

MMSP-1.1: DISTRIBUTED AUDIO-VISUAL PARSING BASED ON MULTIMODAL TRANSFORMER AND DEEP JOINT SOURCE CHANNEL CODING 4623

Penghong Wang, Xiaopeng Fan, Harbin Institute of Technology, China; Jiahui Li, Mengyao Ma, Huawei, China

MMSP-1.2: TALKINGFLOW: TALKING FACIAL LANDMARK GENERATION WITH MULTI-SCALE NORMALIZING FLOW NETWORK 4628

Sen Liang, Zhize Zhou, Hujun Bao, Zhejiang University, China; Rong Li, Zhejiang Lab, China; Juyong Zhang, University of Science and Technology of China, China

MMSP-1.3: INCORPORATING GAZE BEHAVIOR USING JOINT EMBEDDING WITH SCENE CONTEXT FOR DRIVER TAKEOVER DETECTION 4633

Yuning Qiu, Carlos Busso, University of Texas at Dallas, United States of America; Teruhisa Misu, Kumar Akash, Honda Research Institute USA, Inc., United States of America

MMSP-1.4: MULTI-VIEW AND MULTI-MODAL EVENT DETECTION UTILIZING TRANSFORMER-BASED MULTI-SENSOR FUSION 4638

Masahiro Yasuda, Yasunori Ohishi, Shoichiro Saito, Noboru Harada, NTT Corporation, Japan

MMSP-1.5: DISTRIBUTED LABEL DEQUANTIZED GAUSSIAN PROCESS LATENT VARIABLE MODEL FOR MULTI-VIEW DATA INTEGRATION 4643

Koshi Watanabe, Keisuke Maeda, Takahiro Ogawa, Miki Haseyama, Hokkaido University, Japan

MMSP-1.6: CO-ATTENTION-GUIDED BILINEAR MODEL FOR ECHO-BASED DEPTH ESTIMATION 4648

Go Irie, Takashi Shibata, Akisato Kimura, NTT Corporation, Japan

MMSP-2: VIDEO PROCESSING, ANALYSIS, AND SYNTHESIS

MMSP-2.1: MODELING THE DETECTION CAPABILITY OF HIGH-SPEED SPIKING CAMERAS 4653

Junwei Zhao, Zhaoifei Yu, Lei Ma, Ziluo Ding, Shiliang Zhang, Yonghong Tian, Tiejun Huang, Peking University, China

MMSP-2.2: MODERNN: TOWARDS FINE-GRAINED MOTION DETAILS FOR SPATIOTEMPORAL PREDICTIVE LEARNING 4658

Zenghao Chai, Zhengzhuo Xu, Chun Yuan, Tsinghua University, China

MMSP-2.3: GRAPH-BASED POINT CLOUD DENOISING USING SHAPE-AWARE CONSISTENCY FOR FREE-VIEWPOINT VIDEO 4663

Keisuke Nonaka, Ryosuke Watanabe, Haruhisa Kato, KDDI Research, Inc., Japan; Tatsuya Kobayashi, KDDI Research Inc., Japan; Eduardo Pavez, Antonio Ortega, University of Southern California, United States of America

MMSP-2.4: DCSN: DEFORMABLE CONVOLUTIONAL SEMANTIC SEGMENTATION NEURAL NETWORK FOR NON-RIGID SCENES 4668

Bor-Sheng Huang, Chih-Chung Hsu, Han-Yi Kao, Xian-Yun Wang, National Cheng Kung University, Taiwan; Wo-Ting Liao, National Pingtung University of science and technology, Taiwan

MMSP-2.5: TRANSFORMER-BASED DOMAIN ADAPTATION FOR EVENT DATA CLASSIFICATION 4673

Junwei Zhao, Shiliang Zhang, Tiejun Huang, Peking University, China

MMSP-3: EMOTION RECOGNITION

MMSP-3.1: MULTIMODAL EMOTION RECOGNITION WITH SURGICAL AND FABRIC MASKS 4678

Ziqing Yang, Houwei Cao, New York Institute of Technology, United States of America; Katherine Nayan, Rochester Institute of Technology, United States of America; Zehao Fan, New York University Tandon School of Engineering, United States of America

MMSP-3.2: HUMAN EMOTION RECOGNITION USING MULTI-MODAL BIOLOGICAL SIGNALS BASED ON TIME LAG-CONSIDERED CORRELATION MAXIMIZATION 4683

Yuya Moroto, Keisuke Maeda, Takahiro Ogawa, Miki Haseyama, Hokkaido University, Japan

MMSP-3.3: MULTI-MODAL EMOTION RECOGNITION WITH SELF-GUIDED MODALITY CALIBRATION 4688

Mixiao Hou, Zheng Zhang, Guangming Lu, Harbin Institute of Technology, Shenzhen, China

MMSP-3.4: IS CROSS-ATTENTION PREFERABLE TO SELF-ATTENTION FOR MULTI-MODAL EMOTION RECOGNITION? 4693

Vandana Rajan, Andrea Cavallaro, Queen Mary University of London, United Kingdom of Great Britain and Northern Ireland; Alessio Brutti, Fondazione Bruno Kessler, Italy

MMSP-3.5: A PRE-TRAINED AUDIO-VISUAL TRANSFORMER FOR EMOTION RECOGNITION 4698

Minh Tran, Mohammad Soleymani, University of Southern California, United States of America

MMSP-3.6: MEMOBERT: PRE-TRAINING MODEL WITH PROMPT-BASED LEARNING FOR MULTIMODAL EMOTION RECOGNITION 4703

Jinming Zhao, Ruichen Li, Qin Jin, Renmin University of China, China; Xinchao Wang, Haizhou Li, National University of Singapore, Singapore

MMSP-4: OBJECT DETECTION AND TRACKING

MMSP-4.1: GLOBAL-LOCAL FEATURE ENHANCEMENT NETWORK FOR ROBUST OBJECT DETECTION USING MMWAVE RADAR AND CAMERA 4708

Kaikai Deng, Dong Zhao, Qiaoyue Han, Zihan Zhang, Shuyue Wang, Huadong Ma, Beijing University of Posts and Telecommunications, China

MMSP-4.2: LEARNING CORRELATION FOR ONLINE MULTIPLE OBJECT TRACKING 4713

Ying Wang, Chihui Zhuang, Haihui Ye, Yan Yan, Hanzi Wang, Xiamen University, China

MMSP-4.3: BOUNDING BOX DISTRIBUTION LEARNING AND CENTER POINT CALIBRATION FOR ROBUST VISUAL TRACKING 4718

Chihui Zhuang, Yan Yan, Yang Lu, Hanzi Wang, Xiamen University, China; Yanjie Liang, Xiamen University, Peng Cheng Laboratory, China

MMSP-4.4: MULTI-FOCUS GUIDED SEMANTIC AGGREGATION FOR VIDEO OBJECT DETECTION 4723

Haihui Ye, Guangge Wang, Yang Lu, Yan Yan, Hanzi Wang, Xiamen University, China

MMSP-5: MULTIMODAL SIGNAL PROCESSING, ANALYSIS, AND SYNTHESIS II

MMSP-5.1: ENHANCING CONTRASTIVE LEARNING WITH TEMPORAL COGNIZANCE FOR AUDIO-VISUAL REPRESENTATION GENERATION 4728

Chandrashekhar Lavania, Shiva Sundaram, Sundararajan Srinivasan, Katrin Kirchhoff, Amazon, United States of America

MMSP-5.2: CROSS-MODAL KNOWLEDGE DISTILLATION IN MULTI-MODAL FAKE NEWS DETECTION	4733
<i>Zimian Wei, Hengyue Pan, Linbo Qiao, Xin Niu, Peijie Dong, Dongsheng Li, College of Computer, National University of Defense Technology, China</i>	
MMSP-5.3: TRAINING STRATEGIES FOR AUTOMATIC SONG WRITING: A UNIFIED FRAMEWORK PERSPECTIVE	4738
<i>Tao Qian, Shuai Guo, Qin Jin, Renmin University of China, China; Jiatong Shi, Peter Wu, Carnegie Mellon University, United States of America</i>	
MMSP-5.4: RESIDUAL-GUIDED PERSONALIZED SPEECH SYNTHESIS BASED ON FACE IMAGE	4743
<i>Jianrong Wang, Zixuan Wang, Xiaosheng Hu, Xuwei Li, Tianjin University, China; Qiang Fang, Chinese Academy of Social Sciences, China; Li Liu, the Chinese University of Hong Kong, China</i>	
MMSP-5.5: SKETCH STORYTELLING	4748
<i>Yucheng Zhou, Fudan University, China</i>	
MMSP-5.6: MAG+ : AN EXTENDED MULTIMODAL ADAPTATION GATE FOR MULTIMODAL SENTIMENT ANALYSIS	4753
<i>Xianbing Zhao, Yixin Chen, Wanting Li, Buzhou Tang, Harbin Institute of Technology (Shenzhen), China; Lei Gao, University College London, United Kingdom of Great Britain and Northern Ireland</i>	
MMSP-6: MULTIMEDIA DATABASES	
MMSP-6.1: IMAGE-TEXT ALIGNMENT AND RETRIEVAL USING LIGHT-WEIGHT TRANSFORMER	4758
<i>Wenrui Li, Xiaopeng Fan, Harbin Institute of Technology, School of Computer Science & Technology, China</i>	
MMSP-6.2: A GENERAL FRAMEWORK FOR INCOMPLETE CROSS-MODAL RETRIEVAL WITH MISSING LABELS AND MISSING MODALITIES	4763
<i>Mingyang Li, Shao-Lun Huang, Lin Zhang, Tsinghua University, China</i>	
MMSP-6.3: SUBGRAPH REPRESENTATION LEARNING WITH HARD NEGATIVE SAMPLES FOR INDUCTIVE LINK PREDICTION	4768
<i>Heeyoung Kwak, Hyunkyung Bae, Kyomin Jung, Seoul National University, Korea, Republic of</i>	
MMSP-6.4: DEEP HASHING WITH HASH CENTER UPDATE FOR EFFICIENT IMAGE RETRIEVAL	4773
<i>Abin Jose, Daniel Filbert, Christian Rohlfing, Jens-Rainer Ohm, RWTH Aachen University, Germany</i>	
MMSP-6.5: PROTOTYPE-BASED INTER-CAMERA LEARNING FOR PERSON RE-IDENTIFICATION	4778
<i>Lin Wang, Wanqian Zhang, Dayan Wu, Pingting Hong, Bo Li, Institute of Information Engineering, Chinese Academy of Sciences, China</i>	
MMSP-6.6: DHWP: LEARNING HIGH-QUALITY SHORT HASH CODES VIA WEIGHT PRUNING	4783
<i>Zeyu Ma, Guangming Lu, Harbin Institute of Technology, Shenzhen, China; Yuhang Guo, Xiao Luo, Minghua Deng, Peking University, China; Chong Chen, Alibaba Group, China; Wei Cheng, University of Chinese Academy of Sciences, China</i>	
MMSP-7: AUDIO/SPEECH/LANGUAGE ANALYSIS AND SYNTHESIS	
MMSP-7.1: NODE SLICING BROAD LEARNING SYSTEM FOR TEXT CLASSIFICATION	4788
<i>Fagui Liu, Xinjie Wu, Chao Li, South China University of Technology, China</i>	

MMSP-7.2: AUDIO-TEXT RETRIEVAL IN CONTEXT	4793
<i>Siyu Lou, Xuenan Xu, Mengyue Wu, Kai Yu, Shanghai Jiao Tong University, China</i>	
MMSP-7.3: IMPROVED META LEARNING FOR LOW RESOURCE SPEECH RECOGNITION	4798
<i>Satwinder Singh, Ruili Wang, Feng Hou, Massey University, New Zealand</i>	
MMSP-7.4: QUANTIZED WINOGRAD ACCELERATION FOR CONV1D EQUIPPED ASR MODELS ON MOBILE DEVICES	4803
<i>Yiwu Yao, Chengyu Wang, Jun Huang, Alibaba Group, China</i>	
MMSP-7.5: ACOUSTIC-TO-ARTICULATORY INVERSION BASED ON SPEECH DECOMPOSITION AND AUXILIARY FEATURE	4808
<i>Jianrong Wang, Jinyu Liu, Longxuan Zhao, Shanyu Wang, Ruiguo Yu, Tianjin University, China; Li Liu, the Chinese University of Hong Kong, China</i>	
MMSP-7.6: AN AUDIO-SALIENCY MASKING TRANSFORMER FOR AUDIO EMOTION CLASSIFICATION IN MOVIES	4813
<i>Ya-Tse Wu, Jeng-Lin Li, Chi-Chun Lee, National Tsing Hua University, Taiwan</i>	
MMSP-8: JOINT IMAGE AND TEXT PROCESSING	
MMSP-8.1: GENERATIVE ADVERSARIAL NETWORK INCLUDING REFERRING IMAGE SEGMENTATION FOR TEXT-GUIDED IMAGE MANIPULATION	4818
<i>Yuto Watanabe, Ren Togo, Keisuke Maeda, Takahiro Ogawa, Miki Haseyama, Hokkaido University, Japan</i>	
MMSP-8.2: TEXT2POSTER: LAYING OUT STYLIZED TEXTS ON RETRIEVED IMAGES	4823
<i>Chuhao Jin, Hongteng Xu, Ruihua Song, Zhiwu Lu, Renmin University of China, China</i>	
MMSP-8.3: DEEP RANK CROSS-MODAL HASHING WITH SEMANTIC CONSISTENT FOR IMAGE-TEXT RETRIEVAL	4828
<i>Xiaoqing Liu, Huanqiang Zeng, Yifan Shi, Jianqing Zhu, Huaqiao University, China; Kai-Kuang Ma, Nanyang Technological University, China</i>	
MMSP-8.4: VQA-BC: ROBUST VISUAL QUESTION ANSWERING VIA BIDIRECTIONAL CHAINING	4833
<i>Mingrui Lao, Wei Chen, Nan Pu, Michael Lew, Leiden University, Netherlands; Yanming Guo, National University of Defense Technology, China</i>	
MMSP-8.5: TYPE-AWARE MEDICAL VISUAL QUESTION ANSWERING	4838
<i>Anda Zhang, Wei Tao, Ziyang Li, Wenqiang Zhang, Academy for Engineering and Technology, Fudan University, China; Haofen Wang, College of Design and Innovation, Tongji University, China</i>	
MMSP-9: HUMAN CENTRIC MULTIMEDIA	
MMSP-9.1: FROM BOTTOM-UP TO TOP-DOWN: CHARACTERIZATION OF TRAINING PROCESS IN GAZE MODELING	4843
<i>Ron Hecht, Ke Liu, Noa Garnett, Ariel Telpaz, Omer Tsimhoni, General Motors, Israel</i>	
MMSP-9.2: META TALK: LEARNING TO DATA-EFFICIENTLY GENERATE AUDIO-DRIVEN LIP-SYNCHRONIZED TALKING FACE WITH HIGH DEFINITION	4848
<i>Yuhan Zhang, Weihua He, Minglei Li, Kun Tian, Ziyang Zhang, Jie Cheng, Yaoyuan Wang, Jianxing Liao, Huawei Technologies Co., Ltd., China</i>	
MMSP-9.3: MAP: MULTISPECTRAL ADVERSARIAL PATCH TO ATTACK PERSON DETECTION	4853
<i>Taeheon Kim, Hong Joo Lee, Yong Man Ro, KAIST, Korea, Republic of</i>	

MMSP-9.4: GENRE-CONDITIONED LONG-TERM 3D DANCE GENERATION DRIVEN BY MUSIC	4858
<i>Yuhang Huang, Junjie Zhang, Dan Zeng, Shanghai University, China; Shuyan Liu, University of Chinese Academy of Sciences, China; Qian Bao, Wu Liu, JD AI Research, China; Zhineng Chen, Fudan University, China</i>	
MMSP-9.5: LEARNING SOUND LOCALIZATION BETTER FROM SEMANTICALLY SIMILAR SAMPLES	4863
<i>Arda Senocak, Hyeonggon Ryu, In So Kweon, KAIST, Korea, Republic of; Junsik Kim, Harvard University, Korea, Republic of</i>	
MMSP-9.6: BI-DIRECTIONAL MODALITY FUSION NETWORK FOR AUDIO-VISUAL EVENT LOCALIZATION	4868
<i>Shuo Liu, Weize Quan, University of Chinese Academy of Sciences; NLPR, Institute of Automation, Chinese Academy of Sciences, China; Yuan Liu, Speech Lab, Alibaba Group, China; Dong-Ming Yan, NLPR, Institute of Automation, Chinese Academy of Sciences, China</i>	
MMSP-10: IMAGE/GRAPHICS PROCESSING, ANALYSIS, AND CODING	
MMSP-10.1: DYNAMIC MULTI-SCALE LOSS BALANCE FOR OBJECT DETECTION	4873
<i>Yihao Luo, Xiang Cao, Juntao Zhang, Tianjiang Wang, Qi Feng, Huazhong University of Science and Technology, China; Peng Cheng, Coolanyp Limited Liability Company, Kirkland, State of Washington, USA, United States of America</i>	
MMSP-10.2: LATENT SPACE SLICING FOR ENHANCED ENTROPY MODELING IN LEARNING-BASED POINT CLOUD GEOMETRY COMPRESSION	4878
<i>Nicolas Frank, Davi Lazzarotto, Touradj Ebrahimi, École Polytechnique Fédérale de Lausanne (EPFL), Switzerland</i>	
MMSP-10.3: DAM-GAN : IMAGE INPAINTING USING DYNAMIC ATTENTION MAP BASED ON FAKE TEXTURE DETECTION	4883
<i>Dongmin Cha, Daijin Kim, Pohang University of Science and Technology, Korea, Republic of</i>	
MMSP-10.4: IMPROVING REFERENCE-BASED IMAGE COLORIZATION FOR LINE ARTS VIA FEATURE AGGREGATION AND CONTRASTIVE LEARNING	4888
<i>Shukai Wu, Qingqin Wang, Sanyuan Zhang, Zhejiang University, China; Shuchang Xu, Zhejiang University & Hangzhou Normal University & Wenzhou University, China</i>	
MMSP-10.5: FEW-SHOT GAZE ESTIMATION WITH MODEL OFFSET PREDICTORS	4893
<i>Jiawei Ma, Shih-Fu Chang, Columbia University, United States of America; Xu Zhang, Yue Wu, Varsha Hedau, Amazon Alexa AI, United States of America</i>	
MMSP-10.6: ADVERSARIAL EXAMPLES FOR IMAGE CROPPING IN SOCIAL MEDIA	4898
<i>Masatomo Yoshida, Masahiro Okuda, Doshisha University, Japan</i>	
SAM-1: BEAMFORMING	
SAM-1.1: ROBUST ADAPTIVE BEAMFORMING BASED ON POWER METHOD PROCESSING AND SPATIAL SPECTRUM MATCHING	4903
<i>Saeed Mohammadzadeh, Vitor H. Nascimento, University of São Paulo, Brazil; Rodrigo C. de Lamare, Pontifícia Universidade Católica do Rio de Janeiro, Brazil; Osman Kukrer, Eastern Mediterranean University, Cyprus</i>	
SAM-1.2: AN ADAPTIVE ORIENTATIONAL BEAMFORMING TECHNIQUE FOR NARROWBAND INTERFERENCE REJECTION	4908
<i>Jiangyan Han, Boon Poh Ng, Meng Hwa Er, Nanyang Technological University, Singapore</i>	
SAM-1.3: PHASE-ONLY RECONFIGURABLE SPARSE ARRAY BEAMFORMING USING DEEP LEARNING	4913
<i>Syed Ali Hamza, Widener University, United States of America; Moeness Amin, Villanova University, United States of America; Batu Chalise, New York Institute of Technology, United States of America</i>	

SAM-1.4: ROBUST ADAPTIVE BEAMFORMING MAXIMIZING THE WORST-CASE SINR OVER DISTRIBUTIONAL UNCERTAINTY SETS FOR RANDOM INC MATRIX AND SIGNAL STEERING VECTOR	4918
<i>Yongwei Huang, Wenzheng Yang, Guangdong University of Technology, China; Sergiy A. Vorobyov, Aalto University, Finland</i>	
SAM-1.5: IMPROVED BEAMFORMING ENCODING FOR JOINT RADAR AND COMMUNICATION	4923
<i>Tuomas Aittomaki, Visa Koivunen, Aalto University, Finland</i>	
SAM-1.6: STUDY OF THE NULL DIRECTIONS ON THE PERFORMANCE OF DIFFERENTIAL BEAMFORMERS	4928
<i>Xuehan Wang, Jingdong Chen, Northwestern Polytechnical University, China; Israel Cohen, Technion - Israel Institute of Technology, Israel; Jacob Benesty, University of Quebec, Canada</i>	
SAM-2: DOA ESTIMATION AND LEARNING BASED METHODS	
SAM-2.3: DOA M-ESTIMATION USING SPARSE BAYESIAN LEARNING	4933
<i>Christoph Mecklenbraeuer, TU Wien, Austria; Peter Gerstoft, University of California, San Diego, United States of America; Esa Ollila, Aalto University, Finland</i>	
SAM-2.4: LEARNING-AIDED INITIALIZATION FOR VARIATIONAL BAYESIAN DOA ESTIMATION	4938
<i>Yongsung Park, Florian Meyer, Peter Gerstoft, University of California, San Diego, United States of America</i>	
SAM-2.5: NEURAL NETWORK-BASED COMPRESSION FRAMEWORK FOR DOA ESTIMATION EXPLOITING DISTRIBUTED ARRAY	4943
<i>Saidur Pavel, Yimin Zhang, Temple University, United States of America</i>	
SAM-3: WIDEBAND AND MICROPHONE ARRAY PROCESSING	
SAM-3.1: T-SVD BASED BROADBAND NON-SYNCHRONOUS MEASUREMENTS	4948
<i>Long Chen, Weize Sun, Lei Huang, Guitong Chen, Shenzhen University, China</i>	
SAM-3.2: LOW COMPLEX ACCURATE MULTI-SOURCE RTF ESTIMATION	4953
<i>Changheng Li, Jorge Martinez, Richard Christian Hendriks, Delft University of Technology, Netherlands</i>	
SAM-3.3: MULTIPLE OFFSETS MULTILATERATION: A NEW PARADIGM FOR SENSOR NETWORK CALIBRATION WITH UNSYNCHRONIZED REFERENCE NODES	4958
<i>Luca Ferranti, Jani Boutellier, University of Vaasa, Finland; Kalle Åström, Magnus Oskarsson, Lund University, Sweden; Juho Kannala, Aalto University, Finland</i>	
SAM-3.4: REFERENCE MICROPHONE SELECTION AND LOW-RANK APPROXIMATION BASED MULTICHANNEL WIENER FILTER WITH APPLICATION TO SPEECH RECOGNITION	4963
<i>Xing-yu Chen, Jie Zhang, Li-rong Dai, University of Science and Technology of China, China</i>	
SAM-3.5: INCOHERENT SYNTHESIS OF SPARSE BROADBAND ARRAYS BASED ON A PARAMETER-FREE SUBSPACE CLUSTERING	4968
<i>Guy Gubnitsky, University of Haifa, Israel; Yaakov Buchris, Israel Cohen, Technion-Israel Institute of Technology, Israel</i>	
SAM-3.6: INITIALIZATION-FREE IMPLICIT-FOCUSING (IF2) FOR WIDEBAND DIRECTION-OF-ARRIVAL ESTIMATION	4973
<i>Jake Millhiser, Pulak Sarangi, Piya Pal, UC San Diego, United States of America</i>	

SAM-4: MIMO RADAR AND WAVEFORM DESIGN

SAM-4.1: RECURRENT DESIGN OF PROBING WAVEFORM FOR SPARSE BAYESIAN LEARNING BASED DOA ESTIMATION 4978

Linlong Wu, Bhavani Shankar, Ruizhi Hu, Björn Ottersten, University of Luxembourg, Luxembourg; Jisheng Dai, Jiangsu University, China

SAM-4.2: UNIMODULAR WAVEFORM DESIGN WITH LOW CORRELATION LEVELS: A FAST ALGORITHM DEVELOPMENT TO SUPPORT LARGE-SCALE CODE LENGTHS 4983

Yongzhe Li, Chunxuan Shi, Ran Tao, Beijing Institute of Technology, China

SAM-4.3: AIRBORNE MIMO RADAR TRANSMIT-RECEIVE DESIGN UNDER SPECTRAL CONSTRAINT IN SIGNAL-DEPENDENT CLUTTER 4988

Zhihui Li, Junpeng Shi, Dongming Wu, Shujie Shi, Qingsong Zhou, National University of Defense and Technology, China

SAM-4.5: WEAK TARGET DETECTION IN MASSIVE MIMO RADAR VIA AN IMPROVED REINFORCEMENT LEARNING APPROACH 4993

Weitong Zhai, Xiangrong Wang, School of Electronic and Information Engineering, Beihang University, Beijing, 100083, China, China; Maria S. Greco, Fulvio Gini, Department of Information Engineering, University of Pisa, Pisa, Italy, Italy

SAM-4.6: RIS-AIDED MONOSTATIC MIMO RADAR WITH CO-LOCATED ANTENNAS 4998

Stefano Buzzi, Emanuele Grossi, Luca Venturino, University of Cassino and Southern Lazio, Italy; Marco Lops, University of Naples Federico II, Italy

SAM-5: DIRECTION OF ARRIVAL ESTIMATION

SAM-5.1: CONVOLUTIONAL BEAMSPACE USING IIR FILTERS 5003

Po-Chih Chen, P. P. Vaidyanathan, California Institute of Technology, United States of America

SAM-5.2: RATIONAL ARRAYS FOR DOA ESTIMATION 5008

Pranav Kulkarni, P. P. Vaidyanathan, California Institute of Technology, United States of America

SAM-5.3: LOCALIZING MORE SOURCES THAN SENSORS IN PRESENCE OF COHERENT SOURCES 5013

Xinyao Chen, Zai Yang, Xi'an Jiaotong University, China

SAM-5.4: TWO-SNAPSHOT DOA ESTIMATION VIA HANKEL-STRUCTURED MATRIX COMPLETION 5018

Mohammad Bokaei, Saeed Razavikia, Arash Amini, Sharif University of Technology, Iran (Islamic Republic of); Stefano Rini, National Chao Tung University, Taiwan

SAM-5.5: A NOVEL ANGULAR ESTIMATION METHOD IN THE PRESENCE OF NONUNIFORM NOISE 5023

Majdoddin Esfandiari, Sergiy A. Vorobyov, Aalto University, Finland

SAM-5.6: PARTIALLY RELAXED ORTHOGONAL LEAST SQUARES WEIGHTED SUBSPACE FITTING DIRECTION-OF-ARRIVAL ESTIMATION 5028

David Schenck, Katja Lübke, Minh Trinh-Hoang, Marius Pesavento, Technische Universität Darmstadt, Germany

SAM-6: CO-ARRAY BASED ARRAY PROCESSING

SAM-6.1: A NEW COPRIME-ARRAY-BASED CONFIGURATION WITH AUGMENTED DEGREES OF FREEDOM AND REDUCED MUTUAL COUPLING 5033

Nabil Mohsen, Ammar Hawbani, University of Science and Technology of China, China; Monika Agrawal, Indian Institute of Technology Delhi, India; Saeed Alsamhi, Athlone Institute of Technology, Ireland; Liang Zhao, Shenyang Aerospace University, China

SAM-6.2: COARRAY MANIFOLD SEPARATION IN THE SPHERICAL HARMONICS DOMAIN FOR ENHANCED SOURCE LOCALIZATION	5038
<i>Shekhar Kumar Yadav, Nithin V George, Indian Institute of Technology Gandhinagar, India</i>	
SAM-6.3: SPARSE ARRAY SOURCE ENUMERATION VIA COARRAY SUBSPACE OPTIMIZATION	5043
<i>Chun-Lin Liu, National Taiwan University, Taiwan</i>	
SAM-6.4: THE PROTOTYPE CO-PRIME ARRAY WITH A ROBUST DIFFERENCE CO-ARRAY	5048
<i>Ahmed M. A. Shaalan, Jun Du, University of Science and Technology of China, China</i>	
SAM-6.5: DOA ESTIMATION VIA COARRAY TENSOR COMPLETION WITH MISSING SLICES	5053
<i>Hang Zheng, Chengwei Zhou, Zhejiang University, China; André L. F. de Almeida, Federal University of Ceará, Brazil; Yujie Gu, Aptiv, United States of America; Zhiguo Shi, Zhejiang University, International Joint Innovation Center, China</i>	
SAM-6.6: HALF INVERTED NESTED ARRAYS WITH LARGE HOLE-FREE FOURTH-ORDER DIFFERENCE CO-ARRAYS	5058
<i>Yuan-Pon Chen, Chun-Lin Liu, National Taiwan University, Taiwan</i>	
 SAM-7: TRACKING AND LOCALIZATION	
SAM-7.1: SPHERICAL CONVOLUTIONAL RECURRENT NEURAL NETWORK FOR REAL-TIME SOUND SOURCE TRACKING	5063
<i>Tianle Zhong, University of Electronic Science and Technology of China, China; Israel Mendoza Velazquez, Yi Ren, Yoichi Haneda, The University of Electro-Communications, Japan; Hector Manuel Perez Meana, National Polytechnic Institute of Mexico, Mexico</i>	
SAM-7.2: AUDIO-VISUAL TRACKING OF MULTIPLE SPEAKERS VIA A PMBM FILTER	5068
<i>Jinzheng Zhao, Peipei Wu, Xubo Liu, Wenwu Wang, University of Surrey, United Kingdom of Great Britain and Northern Ireland; Yong Xu, Tencent, United States of America; Lyudmila Mihaylova, University of Sheffield, United Kingdom of Great Britain and Northern Ireland; Simon Godsill, University of Cambridge, United Kingdom of Great Britain and Northern Ireland</i>	
SAM-7.3: FLOOR PLAN RECONSTRUCTION WITH HIGH-PRECISION RF-BASED TRACKING	5073
<i>Guozhen Zhu, K. J. Ray Liu, University of Maryland, College Park, United States of America; Chenshu Wu, University of Hong Kong, China; Beibei Wang, Origin Wireless, Inc., United States of America</i>	
SAM-7.4: PARTIAL ARITHMETIC CONSENSUS BASED DISTRIBUTED INTENSITY PARTICLE FLOW SMC-PHD FILTER FOR MULTI-TARGET TRACKING	5078
<i>Peipei Wu, Jinzheng Zhao, Shidrok Goudarzi, Wenwu Wang, University of Surrey, United Kingdom of Great Britain and Northern Ireland</i>	
SAM-7.6: MULTI-MODAL RECURRENT FUSION FOR INDOOR LOCALIZATION	5083
<i>Jianyuan Yu, Virginia Polytechnic Institute and State University, United States of America; Pu Wang, Toshiaki Koike-Akino, Philip Orlik, MITSUBISHI ELECTRIC RESEARCH LABORATORIES, United States of America</i>	
 SAM-8: IMAGING AND ENVIRONMENT MONITORING	
SAM-8.1: IMPROVING JOINT SPARSE HYPERSPECTRAL UNMIXING BY SIMULTANEOUSLY CLUSTERING PIXELS ACCORDING TO THEIR MIXTURES	5088
<i>Seyyede Fatemeh Seyyedsalehi, Hamid Rabiee, Sharif University of Technology, Iran (Islamic Republic of)</i>	

SAM-8.2: GLOBAL EVOLUTION NEURAL NETWORK FOR SEGMENTATION OF REMOTE SENSING IMAGES	5093
<i>Xinzhe Geng, Tao Lei, Qi Chen, Jian Su, Xi He, School of Electronic Information and Artificial Intelligence, Shaanxi University of Science and Technology, China; Qi Wang, School of Computer Science and the Center for OPTical IMagery Analysis and Learning (OPTIMAL), Northwestern Polytechnical University, China; Asoke K. Nandi, Department of Electronic and Computer Engineering, Brunel University London, United Kingdom of Great Britain and Northern Ireland</i>	
SAM-8.3: SPECTRAL-SPATIAL SYMMETRICAL AGGREGATION CROSS-LINKING MULTI-MODAL DATA FUSION NETWORK	5098
<i>Jinping Wang, Jun Li, Xiaojun Tan, Sun Yat-sen University, China</i>	
SAM-8.4: RELATION DISCOVERY IN NONLINEARLY RELATED LARGE-SCALE SETTINGS	5103
<i>Ali Vosoughi, Adora DSouza, Anas Abidin, Axel Wismuller, University of Rochester, United States of America</i>	
SAM-8.5: ACOUSTIC IMAGING ABOARD THE INTERNATIONAL SPACE STATION (ISS): CHALLENGES AND PRELIMINARY RESULTS	5108
<i>Luca Bondi, Charles Shelton, Samarjit Das, Bosch Research, United States of America; Gabriel Chuang, Adarsh Dave, Carnegie Mellon University, United States of America; Christopher Ick, New York University, United States of America; Brian Coltin, KBR, United States of America; Trey Smith, NASA Ames Research Center, United States of America</i>	
SAM-9: SOURCE LOCALIZATION	
SAM-9.1: CONJUGATE AUGMENTED SPATIAL-TEMPORAL NEAR-FIELD SOURCES LOCALIZATION WITH CROSS ARRAY	5113
<i>Zhiwei Jiang, Hua Chen, Ye Tian, Gang Wang, Ningbo University, China; Wei Liu, University of Sheffield, United Kingdom of Great Britain and Northern Ireland</i>	
SAM-9.2: PARAMETRIC MODELS FOR DOA TRAJECTORY LOCALIZATION	5118
<i>Ruchi Pandey, Santosh Nannuru, IIIT Hyderabad, India</i>	
SAM-9.3: JOINT SOURCE LOCALIZATION AND ASSOCIATION THROUGH OVERCOMPLETE REPRESENTATION UNDER MULTIPATH PROPAGATION ENVIRONMENT	5123
<i>Yuan Liu, Zhi-Wei Tan, Andy W. H. Khong, Nanyang Technological University, Singapore; Hongwei Liu, Xidian University, China</i>	
SAM-9.4: SEMIDEFINITE RELAXATION METHOD FOR MOVING OBJECT LOCALIZATION USING A STATIONARY TRANSMITTER AT UNKNOWN POSITION	5128
<i>Ruichao Zheng, Gang Wang, Ningbo University, China; K. C. Ho, University of Missouri, United States of America; Lei Huang, Shenzhen University, China</i>	
SAM-9.5: UNDERDETERMINED TWO-DIMENSIONAL LOCALIZATION FOR WIDEBAND SOURCES BASED ON DISTRIBUTED SENSOR ARRAY NETWORKS	5133
<i>Hantian Wu, Qing Shen, Yibao Liang, Beijing Institute of Technology, China; Wei Liu, University of Sheffield, United Kingdom of Great Britain and Northern Ireland</i>	
SAM-9.6: DIRECT LOCALIZATION: AN ISING MODEL APPROACH	5138
<i>Shiva Akbari, Shahrokh Valaee, University of Toronto, Canada</i>	
SAM-10: DETECTION AND ESTIMATION	
SAM-10.4: TRANSIENT DETECTION WITH UNKNOWN STATISTICS VIA SOURCE CODING	5143
<i>Andrew Finelli, Peter Willett, Yaakov Bar-Shalom, Stefano Marano, University of Connecticut, United States of America</i>	

SAM-10.5: IDENTIFICATION OF PULSE STREAMS OF UNKNOWN SHAPE FROM TIME ENCODING MACHINE SAMPLES	5148
<i>Meghna Kalra, Kiryung Lee, Ohio State University, United States of America; Yoram Bresler, University of Illinois at Urbana-Champaign, United States of America</i>	
SAM-10.6: EXACT SPARSE SUPER-RESOLUTION VIA MODEL AGGREGATION	5153
<i>Hongqing Yu, Heng Qiao, Shanghai Jiao Tong University, China</i>	
SAM-11: ARRAY CALIBRATION AND PERFORMANCE ANALYSIS	
SAM-11.2: A CRLB ANALYSIS OF AOA ESTIMATION USING BLUETOOTH 5	5158
<i>Wentao Shi, Baoqi Huang, Kai Sun, Inner Mongolia University, China</i>	
SAM-11.3: CRAMER-RAO BOUND ANALYSIS OF DISTRIBUTED DOA ESTIMATION EXPLOITING MIXED-PRECISION COVARIANCE MATRIX	5163
<i>Md. Waqeeb T. S. Chowdhury, Yimin Zhang, Temple University, United States of America</i>	
SAM-11.4: CRAMÉR-RAO BOUND AND ANTENNA SELECTION OPTIMIZATION FOR DUAL RADAR-COMMUNICATION DESIGN	5168
<i>Zhaoyi Xu, Athina Petropulu, Rutgers, the State University of New Jersey, United States of America; Fan Liu, Southern University of Science and Technology, China</i>	
SAM-11.5: INFORMATION THEORETIC LIMITS FOR STANDARD AND ONE-BIT COMPRESSED SENSING WITH GRAPH-STRUCTURED SPARSITY	5173
<i>Adarsh Barik, Jean Honorio, Purdue University, United States of America</i>	
SAM-12: MULTICHANNEL PROCESSING AND COMMUNICATIONS	
SAM-12.3: THE DATA/IDENTITY TRADEOFF WITH CENSORED SENSORS	5178
<i>Zachariah Sutton, Peter Willett, University of Connecticut, United States of America; Stefano Marano, University of Salerno, Italy</i>	
SAM-12.4: DOUBLE-RIS VERSUS SINGLE-RIS AIDED SYSTEMS: TENSOR-BASED MIMO CHANNEL ESTIMATION AND DESIGN PERSPECTIVES	5183
<i>Khaled Ardah, Sepideh Gherekhloo, Martin Haardt, Communications Research Laboratory (CRL)/TU Ilmenau, Germany; André L. F. de Almeida, GTEL/Universidade Federal do Ceará (UFC), Brazil</i>	
SAM-12.5: EFFICIENT TWO-STAGE BEAM TRAINING AND CHANNEL ESTIMATION FOR RIS-AIDED MMWAVE SYSTEMS VIA FAST ALTERNATING LEAST SQUARES	5188
<i>Hyeonjin Chung, Sunwoo Kim, Hanyang University, Korea, Republic of</i>	
SPCOM-1: MACHINE LEARNING FOR COMMUNICATION SYSTEMS	
SPCOM-1.1: DEEP JOINT SOURCE-CHANNEL CODING FOR WIRELESS IMAGE TRANSMISSION WITH ADAPTIVE RATE CONTROL	5193
<i>Mingyu Yang, Hun-Seok Kim, University of Michigan, United States of America</i>	
SPCOM-1.2: DEEP SEQUENTIAL BEAMFORMER LEARNING FOR MULTIPATH CHANNELS IN MMWAVE COMMUNICATION SYSTEMS	5198
<i>Aditya Sant, University of California San Diego, United States of America; Afshin Abdi, Joseph Soriaga, Qualcomm Technologies Inc., United States of America</i>	
SPCOM-1.3: DATA-DRIVEN OPTIMIZATION FOR ZERO-DELAY LOSSY SOURCE CODING WITH SIDE INFORMATION	5203
<i>Elad Domanovitz, Daniel Severo, Ashish Khisti, Wei Yu, University of Toronto, Canada</i>	

SPCOM-1.4: DISTRIBUTED IMAGE TRANSMISSION USING DEEP JOINT SOURCE-CHANNEL CODING	5208
<i>Sixian Wang, Ke Yang, Jincheng Dai, Kai Niu, Beijing University of Posts and Telecommunications, China</i>	
SPCOM-1.5: ADAPTIVE WIRELESS POWER ALLOCATION WITH GRAPH NEURAL NETWORKS	5213
<i>Navid NaderiAlizadeh, Alejandro Ribeiro, University of Pennsylvania, United States of America; Mark Eisen, Intel Corporation, United States of America</i>	
SPCOM-1.6: RESTLESS MULTI-ARMED BANDITS UNDER EXOGENOUS GLOBAL MARKOV PROCESS	5218
<i>Tomer Gafni, Ben-Gurion university of the Negev, Israel; Michal Yemini, Princeton University, United States of America; Kobi Cohen, Ben-Gurion University of the Negev, Israel</i>	
SPCOM-2: DECENTRALIZED LEARNING IN COMMUNICATION SYSTEMS	
SPCOM-2.1: BYZANTINE-ROBUST AND COMMUNICATION-EFFICIENT DISTRIBUTED NON-CONVEX LEARNING OVER NON-IID DATA	5223
<i>Xuechao He, Sun Yat-Sen University, China; Heng Zhu, University of Science and Technology of China, China; Qing Ling, Sun Yat-sen University, China</i>	
SPCOM-2.2: COMMUNICATION-EFFICIENT ONLINE FEDERATED LEARNING FRAMEWORK FOR NONLINEAR REGRESSION	5228
<i>Vinay Chakravarthi Gogineni, Stefan Werner, Norwegian University of Science and Technology-NTNU, Norway; Yih-Fang Huang, University of Notre Dame, United States of America; Anthony Kuh, University of Hawaii, United States of America</i>	
SPCOM-2.3: S2 REDUCER: HIGH-PERFORMANCE SPARSE COMMUNICATION TO ACCELERATE DISTRIBUTED DEEP LEARNING	5233
<i>Keshi Ge, Yongquan Fu, Yiming Zhang, Zhiquan Lai, Xiaoge Deng, Dongsheng Li, National University of Defense Technology, China</i>	
SPCOM-2.4: A DATA-DRIVEN QUANTIZATION DESIGN FOR DISTRIBUTED TESTING AGAINST INDEPENDENCE WITH COMMUNICATION CONSTRAINTS	5238
<i>Sebastian Espinosa, Jorge Silva, Universidad de Chile, Chile; Pablo Piantanida, CentraleSupélec CNRS Université Paris-Saclay, France; , ,</i>	
SPCOM-2.5: POWER ALLOCATION FOR WIRELESS FEDERATED LEARNING USING GRAPH NEURAL NETWORKS	5243
<i>Boning Li, Santiago Segarra, Rice University, United States of America; Ananthram Swami, U.S. Army Combat Capabilities Development Command (DEVCOM) Army Research Laboratory, United States of America</i>	
SPCOM-2.6: FEDERATED MULTI-ARMED BANDIT VIA UNCOORDINATED EXPLORATION	5248
<i>Zirui Yan, Quan Xiao, Tianyi Chen, Ali Tajer, Rensselaer Polytechnic institute, United States of America</i>	
SPCOM-2.7: BYZANTINE-RESILIENT DECENTRALIZED COLLABORATIVE LEARNING	5253
<i>Jian Xu, Shao-Lun Huang, Tsinghua University, China</i>	
SPCOM-3: DETECTION AND ESTIMATION FOR COMMUNICATIONS	
SPCOM-3.1: ADAPTIVE IDENTIFICATION OF UNDERWATER ACOUSTIC CHANNEL WITH A MIX OF STATIC AND TIME-VARYING PARAMETERS	5258
<i>Maciej Niedzwiecki, Artur Gancza, Gdansk University of Technology, Poland; Lu Shen, Yuriy Zakharov, University of York, United Kingdom of Great Britain and Northern Ireland</i>	

SPCOM-3.2: ITERATIVE CHANNEL ESTIMATION AND DATA DETECTION	5263
ALGORITHM FOR OTFS MODULATION	
<i>Rabah Ouchikh, Mustapha Djeddou, École Militaire Polytechnique, Algeria; Abdeldjalil Aissa-El-Bey, Thierry Chonavel, École nationale supérieure Mines-Télécom Atlantique Bretagne-Pays de la Loire, France</i>	
SPCOM-3.3: AN ASYMPTOTICALLY OPTIMAL APPROXIMATION OF THE	5268
CONDITIONAL MEAN CHANNEL ESTIMATOR BASED ON GAUSSIAN MIXTURE MODELS	
<i>Michael Koller, Benedikt Fesl, Nurettin Turan, Wolfgang Utschick, Technical University Munich, Germany</i>	
SPCOM-3.4: LOW COMPLEXITY EQUALIZATION FOR AFDM IN DOUBLY	5273
DISPERSIVE CHANNELS	
<i>Ali Bemani, Marios Kountouris, Eurecom, France; Nassar Ksairi, Huawei France R&D, France</i>	
SPCOM-3.5: CSI CLUSTERING WITH VARIATIONAL AUTOENCODING	5278
<i>Michael Baur, Michael Würth, Michael Koller, Vlad-Costin Andrei, Wolfgang Utschick, Technical University of Munich, Germany</i>	
SPCOM-3.6: MASSIVE UNSOURCED RANDOM ACCESS BASED ON BILINEAR	5283
VECTOR APPROXIMATE MESSAGE PASSING	
<i>Ramzi Ayachi, Mohamed Akrouf, Volodymyr Shyianov, Faouzi Bellili, Amine Mezghani, University of Manitoba, Canada</i>	
 SPCOM-4: RESOURCE ALLOCATION IN COMMUNICATION SYSTEMS	
SPCOM-4.1: OPTIMAL QOS-AWARE NETWORK SLICING FOR SERVICE-ORIENTED	5288
NETWORKS WITH FLEXIBLE ROUTING	
<i>Wei-Kun Chen, Beijing Institute of Technology, China; Ya-Feng Liu, Yu-Hong Dai, Academy of Mathematics and Systems Science, Chinese Academy of Sciences, China; Zhi-Quan Luo, Shenzhen Research Institute of Big Data and The Chinese University of Hong Kong, China</i>	
SPCOM-4.2: BYZANTINE-RESILIENT DECENTRALIZED RESOURCE ALLOCATION	5293
<i>Runhua Wang, Qing Ling, Sun Yat-Sen University, China; Yaohua Liu, Nanjing University of Information Science and Technology, China</i>	
SPCOM-4.3: STABILITY ANALYSIS OF UNFOLDED WMMSE FOR POWER	5298
ALLOCATION	
<i>Arindam Chowdhury, Fernando Gama, Santiago Segarra, Rice University, United States of America</i>	
SPCOM-4.4: MONOTONIC GENERALIZED NASH GAMES WITH APPLICATION TO	5303
THE MANAGEMENT OF ENERGY-AWARE ALOHA NETWORKS	
<i>Wenbo Wang, Amir Leshem, Bar-Ilan University, Israel</i>	
SPCOM-4.5: DISTRIBUTED LINK SPARSIFICATION FOR SCALABLE SCHEDULING	5308
USING GRAPH NEURAL NETWORKS	
<i>Zhongyuan Zhao, Santiago Segarra, Rice University, United States of America; Ananthram Swami, US Army's DEVCOM Army Research Laboratory, United States of America</i>	
SPCOM-4.6: A PERFORMANCE ANALYSIS FOR MULTI-RIS-ASSISTED FULL DUPLEX	5313
WIRELESS COMMUNICATION SYSTEM	
<i>Farjam Karim, Bishmita Hazarika, Sandeep Kumar Singh, Keshav Singh, National Sun Yat-sen University, Taiwan</i>	
 SPCOM-5: MIMO SIGNAL PROCESSING FOR WIRELESS SYSTEMS	
SPCOM-5.1: JOINT BEAM SELECTION AND PRECODING BASED ON DIFFERENTIAL	5318
EVOLUTION FOR MILLIMETER-WAVE MASSIVE MIMO SYSTEMS	
<i>Yang Liu, Yancheng Hou, Jiaxuan Wei, Yinghui Zhang, Junxing Zhang, Inner Mongolia University, China; Tiankui Zhang, Beijing University of Posts and Telecommunications, China</i>	

SPCOM-5.2: A NOVEL NEGATIVE L1 PENALTY APPROACH FOR MULTIUSER ONE-BIT MASSIVE MIMO DOWNLINK WITH PSK SIGNALING	5323
<i>Zheyu Wu, Ya-Feng Liu, Yu-Hong Dai, Academy of Mathematics and Systems Science/Chinese Academy of Sciences, China; Bo Jiang, Nanjing Normal University, China</i>	
SPCOM-5.3: A SET-THEORETIC APPROACH TO MIMO DETECTION	5328
<i>Jochen Fink, Renato Cavalcante, Zoran Utkovski, Slawomir Stanczak, Fraunhofer Heinrich-Hertz-Institute, Germany</i>	
SPCOM-5.4: DESIGNING A QAM SIGNAL DETECTOR FOR MASSIVE MIMO SYSTEMS VIA PS-ADMM APPROACH	5333
<i>Quan Zhang, Xuyang Zhao, Jiangtao Wang, Yongchao Wang, Xidian University, China</i>	
SPCOM-5.5: POWER-EFFICIENT HYBRID MIMO RECEIVER WITH TASK-SPECIFIC BEAMFORMING USING LOW-RESOLUTION ADCS	5338
<i>Timur Zirtiloglu, Rabia Tugce Yazicigil, Boston University, United States of America; Nir Shlezinger, Ben-Gurion University, Israel; Yonina Eldar, Weizmann Institute of Science, Israel</i>	
SPCOM-5.6: MIMO DETECTION BY VARIATIONAL POSTERIOR INFERENCE	5343
<i>Junbin Liu, Mingjie Shao, Wing-Kin Ma, The Chinese University of Hong Kong, China</i>	
 SPCOM-6: RECONFIGURABLE INTELLIGENT SURFACES, RELAYS, AND DRONES	
SPCOM-6.1: CONTROLLING SMART PROPAGATION ENVIRONMENTS: LONG-TERM VERSUS SHORT-TERM PHASE SHIFT OPTIMIZATION	5348
<i>Trinh Van Chien, Dinh-Hieu Tran, Hieu Van Nguyen, Symeon Chatzinotas, Björn Ottersten, University of Luxembourg, Luxembourg; Lam Thanh Tu, University of Poitiers, France; Marco Di Renzo, Université Paris-Saclay, CNRS, CentraleSupélec, France</i>	
SPCOM-6.2: DEEP ACTOR-CRITIC FOR CONTINUOUS 3D MOTION CONTROL IN MOBILE RELAY BEAMFORMING NETWORKS	5353
<i>Spilios Evmorfos, Athina Petropulu, RUTGERS UNIVERSITY, United States of America</i>	
SPCOM-6.3: AERIAL BASE STATION PLACEMENT LEVERAGING RADIO TOMOGRAPHIC MAPS	5358
<i>Daniel Romero, Viet Pham, University of Agder, Norway; Geert Leus, TU Delft, Netherlands</i>	
SPCOM-6.5: ATOMIC NORM BASED LOCALIZATION AND ORIENTATION ESTIMATION FOR MILLIMETER-WAVE MIMO OFDM SYSTEMS	5363
<i>Jianxiu Li, Maxime Ferreira Da Costa, Urbashi Mitra, University of Southern California, United States of America</i>	
SPCOM-6.6: ESTIMATION OF CHANNELS IN SYSTEMS WITH INTELLIGENT REFLECTING SURFACES	5368
<i>Michael Joham, Hangze Gao, Wolfgang Utschick, Technical University of Munich, Germany</i>	
 SPCOM-7: CELL-FREE AND DISTRIBUTED ANTENNA SYSTEMS	
SPCOM-7.1: DISTRIBUTED HYBRID BEAMFORMING FOR MMWAVE CELL-FREE MASSIVE MIMO	5373
<i>Nuan Song, Tao Yang, Nokia Bell Labs China, China</i>	
SPCOM-7.2: QUANTIZATION-AWARE PRECODING FOR MU-MIMO WITH LIMITED-CAPACITY FRONTHAUL	5378
<i>Yasaman Khorsandmanesh, Emil Björnson, Joakim Jaldén, KTH Royal Institute of Technology, Sweden</i>	

SPCOM-7.3: AN ONLINE THROUGHPUT MAXIMIZATION ALGORITHM FOR GREEN COORDINATED MULTI-POINT SYSTEMS	5383
<i>Yanjie Dong, Jianqiang Li, Shenzhen University, China; Haijun Zhang, University of Science and Technology Beijing, China; F. Richard Yu, Carleton University, Canada; Song Guo, The Hong Kong Polytechnic University, Hong Kong; Victor C. M. Leung, Shenzhen University/The University of British Columbia, China</i>	
SPCOM-7.4: EFFICIENTLY AND GLOBALLY SOLVING JOINT BEAMFORMING AND COMPRESSION PROBLEM IN THE COOPERATIVE CELLULAR NETWORK VIA LAGRANGIAN DUALITY	5388
<i>Xilai Fan, Ya-Feng Liu, Academy of Mathematics and Systems Science/Chinese Academy of Sciences, China; Liang Liu, The Hong Kong Polytechnic University, China</i>	
SPCOM-7.5: CELL-FREE MASSIVE MIMO: EXPLOITING THE WAX DECOMPOSITION	5393
<i>Juan Vidal Alegría, Fredrik Rusek, Lund University, Sweden; Jinliang Huang, Huawei Technologies AB, Sweden</i>	
SPTM-1: SAMPLING, MULTIRATE, AND DIGITAL SIGNAL PROCESSING I	
SPTM-1.1: LEARNING STRUCTURED SPARSITY FOR TIME-FREQUENCY RECONSTRUCTION	5398
<i>Lei Jiang, Haijian Zhang, Lei Yu, Wuhan University, China</i>	
SPTM-1.2: UNLIMITED SAMPLING WITH SPARSE OUTLIERS: EXPERIMENTS WITH IMPULSIVE AND JUMP OR RESET NOISE	5403
<i>Ayush Bhandari, Imperial College London, United Kingdom of Great Britain and Northern Ireland</i>	
SPTM-1.3: LEARNING APPROACH FOR FAST APPROXIMATE MATRIX FACTORIZATIONS	5408
<i>Haiyan Yu, Zhen Qin, Zhihui Zhu, University of Denver, United States of America</i>	
SPTM-1.4: PARAMETER ESTIMATION IN SPARSE INVERSE PROBLEMS USING BERNOULLI-GAUSSIAN PRIOR	5413
<i>Pierre Barbault, Matthieu Kowalski, Universite Paris Saclay, France; Charles Soussen, CentraleSupélec, France</i>	
SPTM-1.5: SPARSE RECOVERY OF ACOUSTIC WAVES	5418
<i>Mohamed Mansour, Amazon, United States of America</i>	
SPTM-1.6: NONLINEAR SIGNAL DECOMPOSITION BASED ON BLOCK SPARSE APPROXIMATION	5423
<i>El Hadji Diop, University Iba Der Thiam of Thies, Senegal; Karl Skretting, University of Stavanger, Senegal</i>	
SPTM-2: OPTIMIZATION METHODS FOR SIGNAL PROCESSING I	
SPTM-2.1: BLOCK-ACTIVATED ALGORITHMS FOR MULTICOMPONENT FULLY NONSMOOTH MINIMIZATION	5428
<i>Minh N. Bui, Patrick L. Combettes, Zev C. Woodstock, North Carolina State University, United States of America</i>	
SPTM-2.2: BLOCK-COORDINATE FRANK-WOLFE ALGORITHM AND CONVERGENCE ANALYSIS FOR SEMI-RELAXED OPTIMAL TRANSPORT PROBLEM	5433
<i>Takumi Fukunaga, Hiroyuki Kasai, Waseda University, Japan</i>	
SPTM-2.3: AN IMPLICIT GRADIENT-TYPE METHOD FOR LINEARLY CONSTRAINED BILEVEL PROBLEMS	5438
<i>Ioannis Tsaknakis, Prashant Khanduri, Mingyi Hong, University of Minnesota Twin Cities, United States of America</i>	

SPTM-2.4: SCREEN & RELAX: ACCELERATING THE RESOLUTION OF ELASTIC-NET BY SAFE IDENTIFICATION OF THE SOLUTION SUPPORT	5443
<i>Théo Guyard, Univ Rennes, INSA Rennes, CNRS, IRMAR-UMR 6625, F-35000 Rennes, France, France; Cédric Herzet, INRIA Rennes-Bretagne Atlantique, Campus de Beaulieu, 35000 Rennes, France, France; Clément Elvira, SCEE/IETR UMR CNRS 6164, CentraleSupélec, 35510 Cesson Sévigné, France, France</i>	
SPTM-2.5: NODE-SCREENING TESTS FOR THE L0-PENALIZED LEAST-SQUARES PROBLEM	5448
<i>Théo Guyard, Univ Rennes, INSA Rennes, CNRS, IRMAR-UMR 6625, F-35000 Rennes, France, France; Cédric Herzet, INRIA Rennes-Bretagne Atlantique, Campus de Beaulieu, 35000 Rennes, France, France; Clément Elvira, SCEE/IETR UMR CNRS 6164, CentraleSupélec, 35510 Cesson Sévigné, France, France</i>	
SPTM-2.6: PROXIMAL-BASED ADAPTIVE SIMULATED ANNEALING FOR GLOBAL OPTIMIZATION	5453
<i>Thomas Guilmeau, Emilie Chouzenoux, Université Paris-Saclay, CentraleSupélec, INRIA, France; Víctor Elvira, University of Edinburgh, United Kingdom of Great Britain and Northern Ireland</i>	
SPTM-3: GRAPH REPRESENTATION AND ANALYSIS I	
SPTM-3.1: GRAPHON-AIDED JOINT ESTIMATION OF MULTIPLE GRAPHS	5458
<i>Madeline Navarro, Santiago Segarra, William Marsh Rice University, United States of America</i>	
SPTM-3.2: EXPLORING DEEPER GRAPH CONVOLUTIONS FOR SEMI-SUPERVISED NODE CLASSIFICATION	5463
<i>Ashish Tiwari, Shanmuganathan Raman, IIT Gandhinagar, India; Richeek Das, IIT Bombay, India</i>	
SPTM-3.3: DYNAMIC PORTFOLIO CUTS: A SPECTRAL APPROACH TO GRAPH-THEORETIC DIVERSIFICATION	5468
<i>Alvaro Arroyo, Bruno Scalzo, Danilo Mandic, Imperial College London, United Kingdom of Great Britain and Northern Ireland; Ljubisa Stankovic, University of Montenegro, United Kingdom of Great Britain and Northern Ireland</i>	
SPTM-3.4: STABILITY OF NEURAL NETWORKS ON MANIFOLDS TO RELATIVE PERTURBATIONS	5473
<i>Zhiyang Wang, Luana Ruiz, Alejandro Ribeiro, University of Pennsylvania, United States of America</i>	
SPTM-3.5: ADA-STNET: A DYNAMIC ADABOOST SPATIO-TEMPORAL NETWORK FOR TRAFFIC FLOW PREDICTION	5478
<i>Jiawei Sun, Jie Li, Chentao Wu, Zili Tang, Shanghai Jiao Tong University, China; Celimuge Wu, The University of Electro-Communications, Japan</i>	
SPTM-3.6: LABEL PROPAGATION ACROSS GRAPHS: NODE CLASSIFICATION USING GRAPH NEURAL TANGENT KERNELS	5483
<i>Artun Bayer, Arindam Chowdhury, Santiago Segarra, Rice University, United States of America</i>	
SPTM-4: BAYESIAN SIGNAL PROCESSING	
SPTM-4.1: DISTRIBUTED PARTICLE FILTERS FOR STATE TRACKING ON THE STIEFEL MANIFOLD USING TANGENT SPACE STATISTICS	5488
<i>Claudio Bordin, Universidade Federal do ABC, Brazil; Caio de Figueredo, Escola Naval, Brazil; Marcelo Bruno, Instituto Tecnológico de Aeronáutica, Brazil</i>	
SPTM-4.2: HUMAN DECISION MAKING WITH BOUNDED RATIONALITY	5493
<i>Baocheng Geng, University of Alabama at Birmingham, United States of America; Qunwei Li, Ant Group, China; Pramod Varshney, Syracuse University, United States of America</i>	

SPTM-4.3: UNROLLING PARTICLES: UNSUPERVISED LEARNING OF SAMPLING DISTRIBUTIONS	5498
<i>Fernando Gama, Nicolas Zilberstein, Richard Baraniuk, Santiago Segarra, Rice University, United States of America</i>	
SPTM-4.4: SCALABLE DATA ASSOCIATION AND MULTI-TARGET TRACKING UNDER A POISSON MIXTURE MEASUREMENT PROCESS	5503
<i>Qing Li, Jiaming Liang, Simon Godsill, University of Cambridge, United Kingdom of Great Britain and Northern Ireland</i>	
SPTM-4.5: ONLINE LEARNING FOR LATENT YULE-SIMON PROCESSES	5508
<i>Asher Hensley, Petar Djuric, Stony Brook University, United States of America</i>	
 SPTM-5: SAMPLING, MULTIRATE, AND DIGITAL SIGNAL PROCESSING II	
SPTM-5.1: COUNTING THE NUMBER OF DIFFERENT SCALING EXPONENTS IN MULTIVARIATE SCALE-FREE DYNAMICS: CLUSTERING BY BOOTSTRAP IN THE WAVELET DOMAIN	5513
<i>Charles-G�rard Lucas, ENS de Lyon, France; Patrice Abry, CNRS, ENS de Lyon, France; Herwig Wendt, CNRS, Universit� de Toulouse, France; Gustavo Didier, Tulane University, United States of America</i>	
SPTM-5.2: ON THE ACQUISITION OF STATIONARY SIGNALS USING UNIFORM ADCS	5518
<i>Peter Neuhaus, Meik D�rpinghaus, Gerhard Fettweis, Technische Universit�t Dresden, Germany; Nir Shlezinger, Ben-Gurion University of the Negev, Israel; Yonina Eldar, Weizmann Institute of Science, Israel</i>	
SPTM-5.3: DATA-DRIVEN ALGORITHMS FOR GAUSSIAN MEASUREMENT MATRIX DESIGN IN COMPRESSIVE SENSING	5523
<i>Yang Sun, Jonathan Scarlett, National University of Singapore, Singapore</i>	
SPTM-5.4: SCATTERING STATISTICS OF GENERALIZED SPATIAL POISSON POINT PROCESSES	5528
<i>Michael Perlmutter, UCLA, United States of America; Jieqian He, Matthew Hirn, Michigan State University, United States of America</i>	
SPTM-5.5: REGULARIZATION USING DENOISING: EXACT AND ROBUST SIGNAL RECOVERY	5533
<i>Ruturaj Gavaskar, Kunal Chaudhury, Indian Institute of Science, India</i>	
SPTM-5.6: GRAPH-STRUCTURED SPARSE REGULARIZATION VIA CONVEX OPTIMIZATION	5538
<i>Hiroki Kuroda, Daichi Kitahara, Ritsumeikan University, Japan</i>	
 SPTM-6: OPTIMIZATION METHODS FOR SIGNAL PROCESSING II	
SPTM-6.1: DECENTRALIZED BILEVEL OPTIMIZATION FOR PERSONALIZED CLIENT LEARNING	5543
<i>Songtao Lu, Xiaodong Cui, Mark Squillante, Brian Kingsbury, Lior Horesh, IBM Thomas J. Watson Research Center, United States of America</i>	
SPTM-6.3: EXTREME-POINT PURSUIT FOR UNIT-MODULUS OPTIMIZATION	5548
<i>Mingjie Shao, Qi Dai, Wing-Kin Ma, The Chinese University of Hong Kong, China</i>	
SPTM-6.4: GENERALIZED MATCHING PURSUITS FOR THE SPARSE OPTIMIZATION OF SEPARABLE OBJECTIVES	5553
<i>Sebastian Ament, Carla Gomes, Cornell University, United States of America</i>	

SPTM-6.6: DELTA DISTANCING: A LIFTING APPROACH TO LOCALIZING ITEMS FROM USER COMPARISONS	5558
<i>Andrew McRae, Austin Xu, Jihui Jin, Namrata Nadagouda, Nauman Ahad, Peimeng Guan, Mark Davenport, Georgia Institute of Technology, United States of America; Santhosh Karnik, Michigan State University, United States of America</i>	
SPTM-7: GRAPH REPRESENTATION AND ANALYSIS II	
SPTM-7.1: DUAL PATH GRAPH CONVOLUTIONAL NETWORKS	5563
<i>Yunhe Li, University of Montreal, Canada; Yaochen Hu, Yingxue Zhang, Huawei Technologies Canada, Canada</i>	
SPTM-7.2: ON THE STABILITY OF LOW PASS GRAPH FILTER WITH A LARGE NUMBER OF EDGE REWIRES	5568
<i>Hoang-Son Nguyen, Yiran He, Hoi-To Wai, The Chinese University of Hong Kong, Hong Kong</i>	
SPTM-7.3: SPATIO-TEMPORAL GRAPH COMPLEMENTARY SCATTERING NETWORKS	5573
<i>Zida Cheng, Siheng Chen, Ya Zhang, Shanghai Jiao Tong University, China</i>	
SPTM-7.4: CONVOLUTIONAL FILTERING IN SIMPLICIAL COMPLEXES	5578
<i>Elvin Isufi, Maosheng Yang, Delft University of Technology, Netherlands</i>	
SPTM-7.5: ANNIHILATION FILTER APPROACH FOR ESTIMATING GRAPH DYNAMICS FROM DIFFUSION PROCESSES	5583
<i>Arun Venkitaraman, Pascal Frossard, EPFL Switzerland, Switzerland</i>	
SPTM-8: SPARSITY-AWARE PROCESSING	
SPTM-8.1: LEARNING GAUSSIAN GRAPHICAL MODELS WITH DIFFERING PAIRWISE SAMPLE SIZES	5588
<i>Lili Zheng, Genevera Allen, Rice University, United States of America</i>	
SPTM-8.2: R-LOCAL UNLABELED SENSING: IMPROVED ALGORITHM AND APPLICATIONS	5593
<i>Ahmed Ali Abbasi, Shuchin Aeron, Tufts ECE, United States of America; Abiy Tasissa, Tufts Department of Mathematics, United States of America</i>	
SPTM-8.3: FEDERATED OVER-AIR ROBUST SUBSPACE TRACKING FROM MISSING DATA	5598
<i>Praneeth Narayanamurthy, Namrata Vaswani, Aditya Ramamoorthy, Iowa State University, United States of America</i>	
SPTM-8.4: ON CONTINUOUS-DOMAIN INVERSE PROBLEMS WITH SPARSE SUPERPOSITIONS OF DECAYING SINUSOIDS AS SOLUTIONS	5603
<i>Rahul Parhi, Robert Nowak, University of Wisconsin-Madison, United States of America</i>	
SPTM-8.6: MULTIPLICATION-AVOIDING VARIANT OF POWER ITERATION WITH APPLICATIONS	5608
<i>Hongyi Pan, Daa Badawi, Runxuan Miao, Erdem Koyuncu, Ahmet Enis Cetin, Univerity of Illiniois at Chicago, United States of America</i>	
SPTM-9: SAMPLING, MULTIRATE, AND DIGITAL SIGNAL PROCESSING III	
SPTM-9.1: BONA FIDE RIESZ PROJECTIONS FOR DENSITY ESTIMATION	5613
<i>Pol del Aguila Pla, Michael Unser, École polytechnique fédérale de Lausanne, Switzerland</i>	

SPTM-9.2: BLIND MODULO ANALOG-TO-DIGITAL CONVERSION OF VECTOR PROCESSES	5617
<i>Amir Weiss, Gregory W. Wornell, Massachusetts Institute of Technology, United States of America; Everest Huang, MIT Lincoln Laboratory, United States of America; Or Ordentlich, Hebrew University of Jerusalem, Israel</i>	
SPTM-9.3: JOINT RADAR-COMMUNICATIONS PROCESSING FROM A DUAL-BLIND DECONVOLUTION PERSPECTIVE	5622
<i>Edwin Vargas, Roman Jacome, Henry Arguello, UNIVERSIDAD INDUSTRIAL DE SANTANDER, Colombia; Kumar Vijay Mishra, Brian Sadler, United States CDC Army Research Laboratory, United States of America</i>	
SPTM-9.4: FAST MULTISCALE DIFFUSION ON GRAPHS	5627
<i>Sibylle Marcotte, ENS Rennes, France; Amélie Barbe, Rémi Gribonval, Titouan Vayer, Paulo Gonçalves, Université de Lyon, Inria, CNRS, ENSL, LIP, France; Marc Sebban, Université de Lyon, UJM-St-Etienne, CNRS, Lab. Hubert Curien, France; Pierre Borgnat, Université de Lyon, CNRS, ENSL, Lab. de Physique, France</i>	
SPTM-9.5: ADAPTIVE VARIATIONAL NONLINEAR CHIRP MODE DECOMPOSITION	5632
<i>Hao Liang, Xinghao Ding, Xiaotong Tu, Yue Huang, Xiamen University, China; Andreas Jakobsson, Lund University, Sweden</i>	
SPTM-9.6: DIFFERENTIATE-AND-FIRE TIME-ENCODING OF FINITE-RATE-OF-INNOVATION SIGNALS	5637
<i>Abijith Jagannath Kamath, Chandra Sekhar Seelamantula, Indian Institute of Science, India</i>	
SPTM-10: SIGNALS ON GRAPHS AND NETWORKS	
SPTM-10.1: GRAPH LEARNING INFORMATION CRITERION	5642
<i>Koki Yamada, Yuichi Tanaka, Tokyo University of Agriculture and Technology, Japan</i>	
SPTM-10.3: EMBEDDING SIGNALS ON GRAPHS WITH UNBALANCED DIFFUSION EARTH MOVER'S DISTANCE	5647
<i>Alexander Tong, Dennis Shung, Manik Kuchroo, Smita Krishnaswamy, Yale University, United States of America; Guillaume Huguet, Amine Natick, Guillaume Lajoie, Guy Wolf, Université de Montréal, Canada</i>	
SPTM-10.5: MESSAGE PASSING-BASED COOPERATIVE LOCALIZATION WITH EMBEDDED PARTICLE FLOW	5652
<i>Lukas Wielandner, Erik Leitinger, Klaus Witrissal, Graz University of Technology, Austria; Florian Meyer, University of California San Diego, United States of America; Bryan Teague, MIT Lincoln Laboratory, United States of America</i>	
SPTM-10.6: A FRAMEWORK FOR PRIVATE COMMUNICATION WITH SECRET BLOCK STRUCTURE	5657
<i>Maxime Ferreira Da Costa, Urbashi Mitra, University of Southern California, United States of America</i>	
SPTM-11: ADAPTIVE SIGNAL PROCESSING	
SPTM-11.1: LMS AND NLMS ALGORITHMS FOR THE IDENTIFICATION OF IMPULSE RESPONSES WITH INTRINSIC SYMMETRIC OR ANTISYMMETRIC PROPERTIES	5662
<i>Jacob Benesty, INRS-EMT, University of Quebec, Canada; Constantin Paleologu, Silviu Ciochina, University Politehnica of Bucharest, Romania; Eduardo Vinicius Kuhn, Khaled Jamal Bakri, Rui Seara, Federal University of Santa Catarina, Brazil</i>	
SPTM-11.2: DECENTRALIZED LEARNING IN THE PRESENCE OF LOW-RANK NOISE	5667
<i>Roula Nassif, Université Côte d'Azur, France; Virginia Bordignon, Ali Sayed, Ecole Polytechnique Fédérale de Lausanne, Switzerland; Stefan Vlaski, Imperial College London, United Kingdom of Great Britain and Northern Ireland</i>	

SPTM-11.3: ADAPTIVE DIFFUSION WITH COMPRESSED COMMUNICATION..... 5672
Marco Carpentiero, Vincenzo Matta, University of Salerno, Italy; Ali H. Sayed, École Polytechnique Fédérale de Lausanne, Switzerland

SPTM-11.4: JOINT CENTRALITY ESTIMATION AND GRAPH IDENTIFICATION FROM 5677
MIXTURE OF LOW PASS GRAPH SIGNALS
Yiran He, Hoi-To Wai, The Chinese University of Hong Kong, Hong Kong

SPTM-11.5: FAIRNESS-AWARE SELECTIVE SAMPLING ON ATTRIBUTED GRAPHS 5682
Oyku Kose, Yanning Shen, University of California, Irvine, United States of America

SPTM-11.6: A SIMPLE GRAPH NEURAL NETWORK VIA LAYER SNIFFER 5687
Dingyi Zeng, Li Zhou, Wanlong Liu, Hong Qu, Wenyu Chen, University of Electronic Science and Technology of China, China

SPTM-12: CLASSIFICATION AND LEARNING

SPTM-12.1: NEW IMPROVED CRITERION FOR MODEL SELECTION IN SPARSE 5692
HIGH-DIMENSIONAL LINEAR REGRESSION MODELS
Prakash Borpatra Gohain, Magnus Jansson, KTH Royal Institute of Technology, Sweden, Sweden

SPTM-12.2: ON THE USE OF GEODESIC TRIANGLES BETWEEN GAUSSIAN 5697
DISTRIBUTIONS FOR CLASSIFICATION PROBLEMS
Antoine Collas, Chengfang Ren, CentraleSupélec, Université Paris-Saclay, France; Florent Bouchard, CNRS, CentraleSupélec, Université Paris-Saclay, France; Guillaume Ginolhac, Université Savoie Mont Blanc, France; Arnaud Breloy, Université Paris Nanterre, France; Jean-Philippe Ovarlez, Onera, Université Paris-Saclay, France

SPTM-12.3: A NON-CONVEX PROXIMAL APPROACH FOR CENTROID-BASED 5702
CLASSIFICATION
Mewe-Hezoudah Kahanam, Laurent Le Brusquet, Ségolène Martin, Jean-Christophe Pesquet, CentraleSupélec, Université Paris-Saclay, France

SPTM-12.4: EXTENDING THE USE OF MDL FOR HIGH-DIMENSIONAL 5707
PROBLEMS: VARIABLE SELECTION, ROBUST FITTING, AND ADDITIVE MODELING
Zhenyu Wei, Thomas Lee, University of California, Davis, United States of America; Raymond Wong, Texas A&M University, United States of America

SPTM-12.5: CLUSTERING COMPLEX SUBSPACES IN LARGE DIMENSIONS 5712
Roberto Pereira, Xavier Mestre, Centre Tecnològic de Telecomunicacions de Catalunya, Spain; David Gregoratti, Software Radio Systems, Spain

SPTM-12.6: ROBUST CLASSIFICATION WITH FLEXIBLE DISCRIMINANT ANALYSIS 5717
IN HETEROGENEOUS DATA
Pierre Houdouin, Frédéric Pascal, University Paris-Saclay, CentraleSupélec, France; Andrew Wang, University of Cambridge, United Kingdom of Great Britain and Northern Ireland; Matthieu Jonckheere, CNRS, France

SPTM-13: SAMPLING, MULTIRATE, AND DIGITAL SIGNAL PROCESSING IV

SPTM-13.1: RESIDUAL RECOVERY ALGORITHM FOR MODULO SAMPLING 5722
Eyar Azar, Satish Mulleti, Yonina Eldar, Weizmann Institute of Science, Rehovot, Israel, Israel

SPTM-13.2: OPERATOR FORMULATION FOR LINEAR TRANSFORMATIONS AND 5727
SIGNAL ESTIMATION IN THE JOINT SPATIAL-SLEPIAN DOMAIN
Adeem Aslam, University of Engineering and Technology, Pakistan; Zubair Khalid, Lahore University of Management Sciences, Pakistan

SPTM-13.3: SAMPLING SET SELECTION FOR GRAPH SIGNALS UNDER ARBITRARY SIGNAL PRIORS	5732
<i>Junya Hara, Yuichi Tanaka, Tokyo University of Agriculture and Technology, Japan</i>	
SPTM-13.4: DETERMINING JOINT PERIODICITIES IN MULTI-TIME DATA WITH SAMPLING UNCERTAINTIES	5737
<i>David Svedberg, Uppsala University, Sweden; Filip Elvander, KU Leuven, Belgium; Andreas Jakobsson, Lund University, Sweden</i>	
SPTM-13.5: UNLIMITED SAMPLING WITH LOCAL AVERAGES	5742
<i>Dorian Florescu, Ayush Bhandari, Imperial College London, United Kingdom of Great Britain and Northern Ireland</i>	
SPTM-13.6: MODULO EVENT-DRIVEN SAMPLING: SYSTEM IDENTIFICATION AND HARDWARE EXPERIMENTS	5747
<i>Dorian Florescu, Ayush Bhandari, Imperial College London, United Kingdom of Great Britain and Northern Ireland</i>	
SPTM-14: SIGNAL PROCESSING ON GRAPHS AND NETWORKS	
SPTM-15: ESTIMATION THEORY AND METHODS I	
SPTM-15.1: POINT-MASS FILTER WITH DECOMPOSITION OF TRANSIENT DENSITY	5752
<i>Petr Tichavský, Czech Academy of Sciences, Czechia; Ondřej Straka, Jindřich Duník, University of West Bohemia, Czechia</i>	
SPTM-15.2: CRAMER-RAO BOUND FOR THE TIME-VARYING POISSON	5757
<i>Xinhui Rong, Victor Solo, UNSW, Australia</i>	
SPTM-15.3: MODEL SELECTION VIA MISSPECIFIED CRAMER-RAO BOUND MINIMIZATION	5762
<i>Nadav Eliran Rosenthal, Joseph Tabrikian, Ben-Gurion University of the Negev, Israel</i>	
SPTM-15.4: ROBUST PARAMETER ESTIMATION BASED ON THE K-DIVERGENCE	5767
<i>Yair Sorek, Koby Todros, Ben-Gurion University of the Negev, Israel</i>	
SPTM-15.5: A CONVEX FORMULATION FOR THE ROBUST ESTIMATION OF MULTIVARIATE EXPONENTIAL POWER MODELS	5772
<i>Nora Ouzir, Jean-Christophe Pesquet, CentraleSupélec Inria Paris Saclay, France; Frédéric Pascal, CentraleSupélec Paris Saclay, France</i>	
SPTM-15.6: CONDITIONALLY FACTORIZED VARIATIONAL BAYES WITH IMPORTANCE SAMPLING	5777
<i>Runze Gan, Simon Godsill, University of Cambridge, United Kingdom of Great Britain and Northern Ireland</i>	
SPTM-16: DETECTION THEORY AND METHODS	
SPTM-16.1: ON THE FALSE ALARM PROBABILITY OF THE NORMALIZED MATCHED FILTER FOR OFF-GRID TARGET DETECTION	5782
<i>Pierre Develter, Jonathan Bosse, Olivier Rabaste, Jean-Philippe Ovarlez, ONERA, France; Philippe Forster, Université Paris-Saclay, France</i>	
SPTM-16.2: A TWO-STREAM INFORMATION FUSION APPROACH TO ABNORMAL EVENT DETECTION IN VIDEO	5787
<i>Yuxing Yang, Syed Mohsen Naqvi, Newcastle University, United Kingdom of Great Britain and Northern Ireland; Zeyu Fu, University of Oxford, United Kingdom of Great Britain and Northern Ireland</i>	

SPTM-16.3: A TEST FOR CONDITIONAL CORRELATION BETWEEN RANDOM VECTORS BASED ON WEIGHTED U-STATISTICS	5792
<i>Marc Vilà, Jaume Riba, Technical University of Catalonia (UPC), Spain</i>	
SPTM-16.4: SEMI-SUPERVISED STANDARDIZED DETECTION OF PERIODIC SIGNALS WITH APPLICATION TO EXOPLANET DETECTION	5797
<i>Sophia Sulis, Univ. Aix-Marseille, Laboratoire d'Astrophysique de Marseille, CNRS, France; David Mary, Lionel Bigot, Univ. Cote d'Azur, Observatoire de la Cote d'Azur, CNRS, France</i>	
SPTM-16.5: JOINT NORMALITY TEST VIA TWO-DIMENSIONAL PROJECTION	5802
<i>Sara ElBouch, GIPSA-Lab, France; Olivier Michel, Pierre Comon, GIPSA-lab, France</i>	
SPTM-16.6: QUICKEST DETECTION OF COMPOSITE AND NON-STATIONARY CHANGES WITH APPLICATION TO PANDEMIC MONITORING	5807
<i>Yuchen Liang, Venugopal Veeravalli, University of Illinois at Urbana-Champaign, United States of America</i>	
SPTM-17: SIGNAL AND INFORMATION PROCESSING OVER GRAPHS I: STATISTICAL APPROACHES	
SPTM-17.1: A STIMULI-RELEVANT DIRECTED DEPENDENCY INDEX FOR TIME SERIES	5812
<i>Payam Shahsavari Baboukani, Sergios Theodoridis, Jan Østergaard, Aalborg University, Denmark</i>	
SPTM-17.2: JOINT INFERENCE OF MULTIPLE GRAPHS WITH HIDDEN VARIABLES FROM STATIONARY GRAPH SIGNALS	5817
<i>Samuel Rey, Andrei Buciulea, Antonio Garcia Marques, King Juan Carlos University, Spain; Madeline Navarro, Santiago Segarra, Rice University, United States of America</i>	
SPTM-17.3: SPARSE-GROUP LOG-SUM PENALIZED GRAPHICAL MODEL LEARNING FOR TIME SERIES	5822
<i>Jitendra Tugnait, Auburn University, United States of America</i>	
SPTM-17.4: WIDE-SENSE STATIONARITY AND SPECTRAL ESTIMATION FOR GENERALIZED GRAPH SIGNAL	5827
<i>Xingchao Jian, Wee Peng Tay, Nanyang Technological University, Singapore</i>	
SPTM-17.5: BLIND EXTRACTION OF EQUITABLE PARTITIONS FROM GRAPH SIGNALS	5832
<i>Michael Scholkemper, Michael Schaub, RWTH Aachen, Germany</i>	
SPTM-17.6: LEARNING SPARSE GRAPHS WITH A CORE-PERIPHERY STRUCTURE	5837
<i>Sravanthi Gurugubelli, Sundeep Prabhakar Chepuri, Indian Institute of Science, India</i>	
SPTM-18: ADAPTATION AND LEARNING OVER GRAPHS	
SPTM-18.1: OPTIMAL COMBINATION POLICIES FOR ADAPTIVE SOCIAL LEARNING	5842
<i>Ping Hu, Virginia Bordignon, Ali H. Sayed, EPFL, Switzerland; Stefan Vlaski, Imperial College London, United Kingdom of Great Britain and Northern Ireland</i>	
SPTM-18.2: SEISMIC FAULT IDENTIFICATION USING GRAPH HIGH-FREQUENCY COMPONENTS AS INPUT TO GRAPH CONVOLUTIONAL NETWORK	5847
<i>Patitapaban Palo, Aurobinda Routray, IIT Kharagpur, India</i>	
SPTM-18.3: DISTRIBUTED GRAPH LEARNING WITH SMOOTH DATA PRIORS	5852
<i>Isabela Cunha Maia Nobre, Pascal Frossard, EPFL, Switzerland; Mireille El Gheche, Sony AI, Switzerland</i>	

SPTM-18.4: ADVERSPARSE: AN ADVERSARIAL ATTACK FRAMEWORK FOR DEEP SPATIAL-TEMPORAL GRAPH NEURAL NETWORKS	5857
<i>Jiayu Li, Shengmin Jin, Makan Fardad, Reza Zafarani, Syracuse University, United States of America; Tianyun Zhang, Cleveland State University, United States of America</i>	
SPTM-18.5: MULTIMODAL GRAPH SIGNAL DENOISING VIA TWOFOLD GRAPH SMOOTHNESS REGULARIZATION WITH DEEP ALGORITHM UNROLLING	5862
<i>Masatoshi Nagahama, Yuichi Tanaka, Tokyo University of Agriculture and Technology, Japan</i>	
SPTM-18.6: HETEROGENEOUS GRAPH NODE CLASSIFICATION WITH MULTI-HOPS RELATION FEATURES	5867
<i>Xiaolong Xu, Lingjuan Lyu, Hong Jin, Weiqiang Wang, Shuo Jia, ANT GROUP, China</i>	
SPTM-19: ESTIMATION, RECONSTRUCTION, AND DECONVOLUTION	
SPTM-19.1: SIGNAL RECOVERY FROM INCONSISTENT NONLINEAR OBSERVATIONS	5872
<i>Patrick L. Combettes, Zev C. Woodstock, North Carolina State University, United States of America</i>	
SPTM-19.2: PERFECT RECONSTRUCTION OF CLASSES OF NON-BANDLIMITED SIGNALS FROM PROJECTIONS WITH UNKNOWN ANGLES	5877
<i>Renke Wang, Roxana Alexandru, Pier Luigi Dragotti, Imperial College London, United Kingdom of Great Britain and Northern Ireland</i>	
SPTM-19.3: SHORT-AND-SPARSE DECONVOLUTION VIA RANK-ONE CONSTRAINED OPTIMIZATION (ROCO)	5882
<i>Cheng Cheng, Wei Dai, Imperial College London, United Kingdom of Great Britain and Northern Ireland</i>	
SPTM-19.6: BLIND EQUALIZATION OF MOVING AVERAGE CHANNELS OVER GALOIS FIELDS	5887
<i>Arie Yeredor, Tel Aviv University, Israel</i>	
SPTM-20: DETECTION AND TRACKING	
SPTM-20.1: SPARSE SUBSPACE TRACKING IN HIGH DIMENSIONS	5892
<i>Trung Thanh Le, Karim Abed-Meraim, University of Orleans, France; Adel Hafiane, INSA Centre Val de Loire, France; Linh Trung Nguyen, VNU University of Engineering and Technology, Viet Nam</i>	
SPTM-20.2: HOW CAN A COGNITIVE RADAR MASK ITS COGNITION?	5897
<i>Kunal Pattanayak, Vikram Krishnamurthy, Cornell University, United States of America; Christopher Berry, Lockheed Martin Advanced Technologies Laboratory, United States of America</i>	
SPTM-20.3: RTSNET: DEEP LEARNING AIDED KALMAN SMOOTHING	5902
<i>Xiaoyong Ni, Guy Revach, Eidgenössische Technische Hochschule Zürich, Switzerland; Nir Shlezinger, Ben-Gurion University of the Negev, Israel; Ruud J. G. van Sloun, Eindhoven University of Technology, Netherlands; Yonina C. Eldar, Weizmann Institute of Science, Israel</i>	
SPTM-20.5: GENERALIZED AUTOCORRELATION ANALYSIS FOR MULTI-TARGET DETECTION	5907
<i>Ye'ela Shalit, Ran Weber, Asaf Abas, Shay Kreymer, Tamir Bendory, Tel Aviv University, Israel</i>	
SPTM-20.6: APPROXIMATING THE LIKELIHOOD RATIO IN LINEAR-GAUSSIAN STATE-SPACE MODELS FOR CHANGE DETECTION	5912
<i>Kostas Tsampourakis, Víctor Elvira, University of Edinburgh, United Kingdom of Great Britain and Northern Ireland</i>	

SPTM-21: SIGNAL AND INFORMATION PROCESSING OVER GRAPHS II: FILTERING AND TRANSFORMS

SPTM-21.1: LEARNING EXPANDING GRAPHS FOR SIGNAL INTERPOLATION 5917
Bishwadeep Das, Elvin Isufi, TU Delft, Netherlands

SPTM-21.2: HODGELETS: LOCALIZED SPECTRAL REPRESENTATIONS OF FLOWS ON SIMPLICIAL COMPLEXES 5922
T. Mitchell Roddenberry, Santiago Segarra, Rice University, United States of America; Florian Frantzen, Michael T. Schaub, RWTH Aachen University, Germany

SPTM-21.3: RECOVERY OF GRAPH SIGNALS FROM SIGN MEASUREMENTS 5927
Wenwei Liu, Hui Feng, Kaixuan Wang, Bo Hu, Fudan University, China; Feng Ji, Nanyang Technological University, Singapore

SPTM-21.4: EDGE SAMPLING OF GRAPHS BASED ON EDGE SMOOTHNESS 5932
Kenta Yanagiya, Koki Yamada, Yuichi Tanaka, Tokyo University of Agriculture and Technology, Japan; Yasuo Katsuhara, Tomoya Takatani, Frontier Research Center, Toyota Motor Corporation, Japan

SPTM-21.5: WLS DESIGN OF ARMA GRAPH FILTERS USING ITERATIVE SECOND-ORDER CONE PROGRAMMING 5937
Darukeesan Pakiyarajah, University of Jaffna, Sri Lanka; Chamira Edussooriya, University of Moratuwa, Sri Lanka

SPTM-21.6: LINEAR-TIME SAMPLING ON SIGNED GRAPHS VIA GERSHGORIN DISC PERFECT ALIGNMENT 5942
Chinthaka Dinesh, Ivan Bajic, Simon Fraser University, Canada; Saghar Bagheri, Gene Cheung, York University, Canada

SPTM-22: SIGNAL PROCESSING OVER NETWORKS

SPTM-22.1: PRIVACY-PRESERVING FEDERATED MULTI-TASK LINEAR REGRESSION: A ONE-SHOT LINEAR MIXING APPROACH INSPIRED BY GRAPH REGULARIZATION 5947
Harlin Lee, Andrea Bertozzi, University of California, Los Angeles, United States of America; Jelena Kovačević, New York University, United States of America; Yuejie Chi, Carnegie Mellon University, United States of America

SPTM-22.2: ECO-FEDSPLIT: FEDERATED LEARNING WITH ERROR-COMPENSATED COMPRESSION 5952
Sarit Khirirat, Division of Decision and Control Systems, KTH Royal Institute of Technology, Sweden, Sweden; Sindri Magnússon, Stockholm University, Sweden; Mikael Johansson, KTH Royal Institute of Technology, Sweden

SPTM-22.3: A TIME ENCODING APPROACH TO TRAINING SPIKING NEURAL NETWORKS 5957
Karen Adam, Ecole Polytechnique Fédérale de Lausanne, Switzerland

SPTM-22.4: TRANSIENT ANALYSIS OF CLUSTERED MULTITASK DIFFUSION RLS ALGORITHM 5962
Wei Gao, Jiangsu University, China; Jie Chen, Wentao Shi, Qunfei Zhang, Northwestern Polytechnical University, China; Cedric Richard, Université Côte d'Azur, France

SPTM-22.5: IMPROVING INFERENCE FOR SPATIAL SIGNALS BY CONTEXTUAL FALSE DISCOVERY RATES 5967
Martin Gözl, Abdelhak M. Zoubir, Technische Universität Darmstadt, Germany; Visa Koivunen, Aalto University, Finland

SPTM-22.6: ESTIMATION OF THE ADMITTANCE MATRIX IN POWER SYSTEMS UNDER LAPLACIAN AND PHYSICAL CONSTRAINTS 5972
Morad Halihal, Tirza Routtenberg, Ben Gurion University of the Negev, Israel

SPTM-23: ESTIMATION THEORY AND APPLICATIONS

SPTM-23.2: INCIPIENT FAULT SEVERITY ESTIMATION USING LOCAL MAHALANOBIS DISTANCE 5977

Junjie Yang, Claude Delpha, Universite Paris Saclay, France

SPTM-23.3: GRIDLESS DOA ESTIMATION UNDER THE MULTI-FREQUENCY MODEL 5982

Yifan Wu, Peter Gerstoft, University of California, San Diego, United States of America; Michael Wakin, Colorado School of Mines, United States of America

SPTM-24: NMF, DIVERGENCES AND OPTIMIZATION

SPTM-24.1: ORTHOGONAL NONNEGATIVE MATRIX TRI-FACTORIZATION FOR COMMUNITY DETECTION IN MULTIPLEX NETWORKS 5987

Meiby Ortiz-Bouza, Selin Aiyente, Michigan State University, United States of America

SPTM-24.2: STUDYING THREE FAMILIES OF DIVERGENCES TO COMPARE WIDE-SENSE STATIONARY GAUSSIAN ARMA PROCESSES 5992

Eric Grivel, Bordeaux INP, France

SPTM-24.3: MULTIVARIATE MULTISCALE COSINE SIMILARITY ENTROPY 5997

Hongjian Xiao, Danilo Mandic, Imperial College London, United Kingdom of Great Britain and Northern Ireland; Theerasak Chanwimalueang, Srinakharinwirot University, Thailand

SPTM-24.4: ZERO-TH-ORDER RANDOMIZED SUBSPACE NEWTON METHODS 6002

Erik Berglund, Sarit Khirirat, Xiaoyu Wang, KTH Royal Institute of Technology, Sweden

SPTM-24.5: FAST AND STABLE CONVERGENCE OF ONLINE SGD FOR CV@R-BASED RISK-AWARE LEARNING 6007

Dionysios Kalogerias, Yale University, United States of America

SPTM-24.6: DEEP INITIALIZATION FOR GUARANTEED UNIMODULAR QUADRATIC PROGRAMMING 6012

Amrutha Varshini Ramesh, University of British Columbia, Canada; Mojtaba Soltanalian, University of Illinois at Chicago, United States of America

SPE-1: SPEECH SEPARATION

SPE-1.1: CONTINUOUS SPEECH SEPARATION WITH RECURRENT SELECTIVE ATTENTION NETWORK 6017

Yixuan Zhang, The Ohio State University, United States of America; Zhuo Chen, Jian Wu, Takuya Yoshioka, Peidong Wang, Zhong Meng, Jinyu Li, Microsoft, United States of America

SPE-1.2: SA-SDR: A NOVEL LOSS FUNCTION FOR SEPARATION OF MEETING STYLE DATA 6022

Thilo von Neumann, Christoph Boeddeker, Reinhold Haeb-Umbach, Paderborn University, Germany; Keisuke Kinoshita, Marc Delcroix, NTT Corporation, Japan

SPE-1.3: VARARRAY: ARRAY-GEOMETRY-AGNOSTIC CONTINUOUS SPEECH SEPARATION 6027

Takuya Yoshioka, Xiaofei Wang, Dongmei Wang, Min Tang, Zirun Zhu, Zhuo Chen, Naoyuki Kanda, Microsoft, United States of America

SPE-1.4: ALL-NEURAL BEAMFORMER FOR CONTINUOUS SPEECH SEPARATION 6032

Zhuohuang Zhang, Indiana University Bloomington, United States of America; Takuya Yoshioka, Naoyuki Kanda, Zhuo Chen, Xiaofei Wang, Dongmei Wang, Sefik Emre Eskimez, Microsoft, United States of America

SPE-1.6: MINING HARD SAMPLES LOCALLY AND GLOBALLY FOR IMPROVED SPEECH SEPARATION	6037
<i>Kai Wang, Yizhou Peng, Hao Huang, Ying Hu, Xinjiang University, China; Sheng Li, National Institute of Information and Communications Technology, China</i>	
SPE-2: SPEECH RECOGNITION: ROBUST SPEECH RECOGNITION I	
SPE-2.1: AUDIO-VISUAL MULTI-CHANNEL SPEECH SEPARATION, DEREVERBERATION AND RECOGNITION	6042
<i>Guinan Li, Jiajun Deng, Xunying Liu, Helen Meng, The Chinese University of Hong Kong, China; Jianwei Yu, Tencent AI lab, China</i>	
SPE-2.2: BEST OF BOTH WORLDS: MULTI-TASK AUDIO-VISUAL AUTOMATIC SPEECH RECOGNITION AND ACTIVE SPEAKER DETECTION	6047
<i>Otavio Braga, Olivier Siohan, Google, United States of America</i>	
SPE-2.3: END-TO-END MULTI-MODAL SPEECH RECOGNITION WITH AIR AND BONE CONDUCTED SPEECH	6052
<i>Junqi Chen, Mou Wang, Northwestern Polytechnical University, China; Xiao-Lei Zhang, Northwestern Polytechnic University, China; Zhiyong Huang, Susanto Rahardja, National University of Singapore, Singapore</i>	
SPE-2.4: END-TO-END SPEECH RECOGNITION WITH JOINT DEREVERBERATION OF SUB-BAND AUTOREGRESSIVE ENVELOPES	6057
<i>Rohit Kumar, Anurenjan Purushothaman, Sriram Ganapathy, Indian Institute of Science, Bangalore, India; Anirudh Sreeram, University of Southern California, United States of America</i>	
SPE-2.5: IMPROVING NOISE ROBUSTNESS OF CONTRASTIVE SPEECH REPRESENTATION LEARNING WITH SPEECH RECONSTRUCTION	6062
<i>Heming Wang, DeLiang Wang, The Ohio State University, United States of America; Yao Qian, Xiaofei Wang, Yiming Wang, Chengyi Wang, Shujie Liu, Takuya Yoshioka, Jinyu Li, Microsoft Corporation, United States of America</i>	
SPE-2.6: MULTI-CHANNEL MULTI-SPEAKER ASR USING 3D SPATIAL FEATURE	6067
<i>Yiwen Shao, Johns Hopkins University, United States of America; Shi-Xiong Zhang, Dong Yu, Tencent AI Lab, United States of America</i>	
SPE-3: SPEECH SYNTHESIS: GENERAL TOPICS I	
SPE-3.1: IMPROVING CROSS-LINGUAL SPEECH SYNTHESIS WITH TRIPLET TRAINING SCHEME	6072
<i>Jianhao Ye, Hongbin Zhou, Zhiba Su, Wendi He, Kaimeng Ren, Lin Li, Heng Lu, Ximalaya Inc., China</i>	
SPE-3.2: IMPROVING PHONETIC REALIZATIONS IN TTS BY USING PHONEME-ALIGNED GRAPHEMES	6077
<i>Manish Sharma, Yizhi Hong, Emily Kaplan, Siamak Tazari, Rob Clark, Google, United Kingdom of Great Britain and Northern Ireland</i>	
SPE-3.3: CONTEXT-AWARE MASK PREDICTION NETWORK FOR END-TO-END TEXT-BASED SPEECH EDITING	6082
<i>Tao Wang, Jiangyan Yi, Ruibo Fu, Jianhua Tao, Zhengqi Wen, Institute of Automation, Chinese Academy of Sciences, China; Liqun Deng, Huawei Noah's Ark Lab, Shenzhen, China, China</i>	
SPE-3.4: A STUDY ON THE EFFICACY OF MODEL PRE-TRAINING IN DEVELOPING NEURAL TEXT-TO-SPEECH SYSTEM	6087
<i>Guangyan Zhang, Daxin Tan, Tan Lee, Department of Electronic Engineering, The Chinese University of Hong Kong, Hong Kong; Yichong Leng, University of Science and Technology of China, China; Ying Qin, Institute of Information Science, Beijing Jiaotong University, China; Kaitao Song, Xu Tan, Microsoft Research Asia, China; Sheng Zhao, Microsoft Azure Speech, China</i>	

SPE-3.6: ONE TTS ALIGNMENT TO RULE THEM ALL	6092
<i>Rohan Badlani, Adrian Lancucki, Kevin J. Shih, Rafael Valle, Wei Ping, Bryan Catanzaro, NVIDIA, United States of America</i>	

SPE-4: LANGUAGE MODELING

SPE-4.1: CAPITALIZATION NORMALIZATION FOR LANGUAGE MODELING WITH AN ACCURATE AND EFFICIENT HIERARCHICAL RNN MODEL	6097
--	-------------

Hao Zhang, You-Chi Cheng, Shankar Kumar, W. Ronny Huang, Mingqing Chen, Rajiv Mathews, Google, United States of America

SPE-4.2: ENHANCE RNNLMS WITH HIERARCHICAL MULTI-TASK LEARNING FOR ASR	6102
--	-------------

Minguang Song, Yunxin Zhao, University of Missouri-Columbia, United States of America

SPE-4.3: NEURAL-FST CLASS LANGUAGE MODEL FOR END-TO-END SPEECH RECOGNITION	6107
---	-------------

Antoine Bruguier, Duc Le, Dangna Li, Zhe Liu, Bo Wang, Eun Chang, Fuchun Peng, Ozlem Kalinli, Mike Seltzer, Facebook, United States of America; Rohit Prabhavalkar, Google, United States of America

SPE-4.4: LATTICEBART: LATTICE-TO-LATTICE PRE-TRAINING FOR SPEECH RECOGNITION	6112
---	-------------

Lingfeng Dai, Lu Chen, Zhikai Zhou, Kai Yu, Shanghai Jiao Tong University, China

SPE-4.5: RESCOREBERT: DISCRIMINATIVE SPEECH RECOGNITION RESCORING WITH BERT	6117
--	-------------

Liyen Xu, Yile Gu, Jari Kolehmainen, Haidar Khan, Ankur Gandhe, Ariya Rastrow, Andreas Stolcke, Ivan Bulyko, Amazon Alexa AI, United States of America

SPE-4.6: HYBRID SUB-WORD SEGMENTATION FOR HANDLING LONG TAIL IN MORPHOLOGICALLY RICH LOW RESOURCE LANGUAGES	6122
--	-------------

Sreeja Manghat, Sreeram Manghat, IEEE Graduate Member, India; Tanja Schultz, Cognitive Systems Lab, University Bremen, Germany, Germany

SPE-5: SPEAKER RECOGNITION I: SELF SUPERVISION

SPE-5.1: SELF-SUPERVISED SPEAKER VERIFICATION WITH SIMPLE SIAMESE NETWORK AND SELF-SUPERVISED REGULARIZATION	6127
---	-------------

Mufan Sang, The University of Texas at Dallas, United States of America; Haoqi Li, Fang Liu, Andrew O. Arnold, Li Wan, Amazon, United States of America

SPE-5.2: SELF-SUPERVISED SPEAKER RECOGNITION TRAINING USING HUMAN-MACHINE DIALOGUES	6132
--	-------------

Metehan Cekic, Upamanyu Madhow, UCSB, United States of America; Ruirui Li, Zeya Chen, Yuguang Yang, Andreas Stolcke, Amazon, United States of America

SPE-5.3: MULTI-TASK VOICE ACTIVATED FRAMEWORK USING SELF-SUPERVISED LEARNING	6137
---	-------------

Shehzeen Hussain, University of California, San Diego, United States of America; Van Nguyen, Shuhua Zhang, Erik Visser, Qualcomm Technologies, Inc., United States of America

SPE-5.4: SELF-SUPERVISED SPEAKER RECOGNITION WITH LOSS-GATED LEARNING	6142
--	-------------

*Ruijie Tao, Ville Hautamäki, Haizhou Li, National University of Singapore, Singapore; Kong Aik Lee, A*STAR, Singapore; Rohan Kumar Das, Fortemedia Singapore, Singapore*

SPE-5.5: LARGE-SCALE SELF-SUPERVISED SPEECH REPRESENTATION LEARNING FOR AUTOMATIC SPEAKER VERIFICATION	6147
<i>Zhengyang Chen, Yanmin Qian, Shanghai Jiao Tong University, China; Sanyuan Chen, Yu Wu, Yao Qian, Chengyi Wang, Shujie Liu, Michael Zeng, Microsoft Corporation, China</i>	
SPE-5.6: UNISPEECH-SAT: UNIVERSAL SPEECH REPRESENTATION LEARNING WITH SPEAKER AWARE PRE-TRAINING	6152
<i>Sanyuan Chen, Xiangzhan Yu, Harbin Institute of Technology, China; Yu Wu, Chengyi Wang, Zhengyang Chen, Zhuo Chen, Shujie Liu, Jian Wu, Yao Qian, Furu Wei, Jinyu Li, Microsoft Corporation, China</i>	
SPE-6: SPEECH AND SPOKEN LANGUAGE CORPORA	
SPE-6.1: CARINA – A CORPUS OF ALIGNED GERMAN READ SPEECH INCLUDING ANNOTATIONS	6157
<i>Hannes Kath, Simon Stone, Stefan Rapp, Peter Birkholz, Technische Universität Dresden, Germany</i>	
SPE-6.2: TOWARDS MEASURING FAIRNESS IN SPEECH RECOGNITION: CASUAL CONVERSATIONS DATASET TRANSCRIPTIONS	6162
<i>Chunxi Liu, Michael Picheny, Leda Sari, Pooja Chitkara, Alex Xiao, Xiaohui Zhang, Mark Chou, Andrés Alvarado, Caner Hazirbas, Yatharth Saraf, Meta AI, United States of America</i>	
SPE-6.3: M2MET: THE ICASSP 2022 MULTI-CHANNEL MULTI-PARTY MEETING TRANSCRIPTION CHALLENGE	6167
<i>Fan Yu, Shiliang Zhang, Siqi Zheng, Zhihao Du, Weilong Huang, Zhijie Yan, Bin Ma, Alibaba, China; Yihui Fu, Lei Xie, Pengcheng Guo, AISHELL Foundation, China; Xin Xu, Hui Bu, Beijing Shell Shell Technology Co., Ltd., China</i>	
SPE-6.4: ADIMA: ABUSE DETECTION IN MULTILINGUAL AUDIO	6172
<i>Vikram Gupta, Rini Sharon, Ramit Sawhney, Debdoot Mukherjee, ShareChat, India, India</i>	
SPE-6.5: ANNO-MI: A DATASET OF EXPERT-ANNOTATED COUNSELLING DIALOGUES	6177
<i>Zixiu Wu, Philips Research and University of Cagliari, Netherlands; Simone Balloccu, Ehud Reiter, University of Aberdeen, United Kingdom of Great Britain and Northern Ireland; Vivek Kumar, Diego Reforgiato Recupero, Daniele Riboni, University of Cagliari, Italy; Rim Helaoui, Philips Research, Netherlands</i>	
SPE-6.6: WENETSPEECH: A 10000+ HOURS MULTI-DOMAIN MANDARIN CORPUS FOR SPEECH RECOGNITION	6182
<i>Binbin Zhang, Northwestern Polytechnical University; Mobvoi Inc, China; Hang Lv, Pengcheng Guo, Qijie Shao, Lei Xie, Northwestern Polytechnical University, China; Chao Yang, Xiaoyu Chen, Chenchen Zeng, Di Wu, Zhendong Peng, Mobvoi Inc., China; Xin Xu, Hui Bu, Beijing Shell Shell Technology Co., Ltd., China</i>	
SPE-7: SPEECH SYNTHESIS: GENERAL TOPICS II	
SPE-7.1: WAVEBENDER GAN: AN ARCHITECTURE FOR PHONETICALLY MEANINGFUL SPEECH MANIPULATION	6187
<i>Gustavo Teodoro Döhler Beck, Ulme Wennberg, Zofia Malisz, Gustav Eje Henter, KTH Royal Institute of Technology, Sweden</i>	
SPE-7.2: FRE-GAN 2: FAST AND EFFICIENT FREQUENCY-CONSISTENT AUDIO SYNTHESIS	6192
<i>Sang-Hoon Lee, Ji-Hoon Kim, Kang-Eun Lee, Seong-Whan Lee, Korea University, Korea, Republic of</i>	
SPE-7.3: R-G2P: EVALUATING AND ENHANCING ROBUSTNESS OF GRAPHEME TO PHONEME CONVERSION BY CONTROLLED NOISE INTRODUCING AND CONTEXTUAL INFORMATION INCORPORATION	6197
<i>Chendong Zhao, Haoqian Wang, The Shenzhen International Graduate School, Tsinghua University, China, China; Jianzong Wang, Xiaoyang Qu, Jing Xiao, Ping An Technology (Shenzhen) Co., Ltd., China</i>	

SPE-7.4: NEURAL GRAPHEME-TO-PHONEME CONVERSION WITH PRE-TRAINED GRAPHEME MODELS	6202
<i>Lu Dong, Zhi-Qiang Guo, Chao-Hong Tan, Zhen-Hua Ling, National Engineering Laboratory for Speech and Language Information Processing, University of Science and Technology of China, Hefei, P. R. China, China; Ya-Jun Hu, Yuan Jiang, iFLYTEK Research, iFLYTEK Co., Ltd., Hefei, P. R. China, China</i>	
SPE-7.5: ISTFTNET: FAST AND LIGHTWEIGHT MEL-SPECTROGRAM VOCODER INCORPORATING INVERSE SHORT-TIME FOURIER TRANSFORM	6207
<i>Takuhiro Kaneko, Kou Tanaka, Hirokazu Kameoka, Shogo Seki, NTT Corporation, Japan</i>	
SPE-7.6: ACOUSTIC APPLICATION OF PHASE RECONSTRUCTION ALGORITHMS IN OPTICS	6212
<i>Tomoki Kobayashi, Tomoro Tanaka, Kohei Yatabe, Yasuhiro Oikawa, Waseda University, Japan</i>	
SPE-8: MACHINE TRANSLATION FOR SPOKEN AND WRITTEN LANGUAGE	
SPE-8.1: CPT: CROSS-MODAL PREFIX-TUNING FOR SPEECH-TO-TEXT TRANSLATION	6217
<i>Yukun Ma, Trung Hieu Nguyen, Bin Ma, Alibaba Group, Singapore</i>	
SPE-8.2: TACKLING DATA SCARCITY IN SPEECH TRANSLATION USING ZERO-SHOT MULTILINGUAL MACHINE TRANSLATION TECHNIQUES	6222
<i>Tu Anh Dinh, Danni Liu, Jan Niehues, Maastricht University, Netherlands</i>	
SPE-8.3: IMPROVING END-TO-END SPEECH TRANSLATION MODEL WITH BERT-BASED CONTEXTUAL INFORMATION	6227
<i>Jeong-Uk Bang, Min-Kyu Lee, Seung Yun, Sang-Hun Kim, Electronics and Telecommunications Research Institute (ETRI), Korea, Republic of</i>	
SPE-8.4: CONTEXT-ADAPTIVE DOCUMENT-LEVEL NEURAL MACHINE TRANSLATION	6232
<i>Linlin Zhang, Zhejiang University, China; Zhirui Zhang, Boxing Chen, Weihua Luo, Luo Si, Alibaba DAMO Academy, China</i>	
SPE-8.5: INTEGRATING MULTIPLE ASR SYSTEMS INTO NLP BACKEND WITH ATTENTION FUSION	6237
<i>Takatomo Kano, Atsunori Ogawa, Marc Delcroix, NTT Corporation, Japan; Shinji Watanabe, Language Technologies Institute, Carnegie Mellon University, United States of America</i>	
SPE-8.6: ISOMETRIC MT: NEURAL MACHINE TRANSLATION FOR AUTOMATIC DUBBING	6242
<i>Sorafel Lakew, Yogesh Virkar, Prashant Mathur, Marcello Federico, Amazon AI, United States of America</i>	
SPE-9: DEPRESSION DETECTION	
SPE-9.1: AUTOMATIC DEPRESSION DETECTION: AN EMOTIONAL AUDIO-TEXTUAL CORPUS AND A GRU/BILSTM-BASED MODEL	6247
<i>Ying Shen, Huiyu Yang, Lin Lin, Tongji University, China</i>	
SPE-9.2: MULTIMODAL DEPRESSION CLASSIFICATION USING ARTICULATORY COORDINATION FEATURES AND HIERARCHICAL ATTENTION BASED TEXT EMBEDDINGS	6252
<i>Nadee Seneviratne, Carol Espy-Wilson, University of Maryland - College Park, United States of America</i>	
SPE-9.3: THIN SLICES OF DEPRESSION: IMPROVING DEPRESSION DETECTION PERFORMANCE THROUGH DATA SEGMENTATION	6257
<i>Rawan Alsarrani, Alessandro Vinciarelli, University of Glasgow, United Kingdom of Great Britain and Northern Ireland; Anna Esposito, Universita' degli Studi della Campania, Italy</i>	

SPE-9.4: CLIMATE AND WEATHER: INSPECTING DEPRESSION DETECTION VIA EMOTION RECOGNITION	6262
<i>Wen Wu, University of Cambridge, United Kingdom of Great Britain and Northern Ireland; Mengyue Wu, Kai Yu, Shanghai Jiao Tong University, China</i>	
SPE-9.5: FRAUG: A FRAME RATE BASED DATA AUGMENTATION METHOD FOR DEPRESSION DETECTION FROM SPEECH SIGNALS	6267
<i>Vijay Ravi, Jinhan Wang, Jonathan Flint, Abeer Alwan, University of California Los Angeles, United States of America</i>	
SPE-9.6: PRIVACY SENSITIVE SPEECH ANALYSIS USING FEDERATED LEARNING TO ASSESS DEPRESSION	6272
<i>Suhas Bn, Saeed Abdullah, Pennsylvania State University, United States of America</i>	
 SPE-10: SPEECH RECOGNITION: ROBUST SPEECH RECOGNITION II	
SPE-10.1: A TIME DOMAIN PROGRESSIVE LEARNING APPROACH WITH SNR CONSTRICTION FOR SINGLE-CHANNEL SPEECH ENHANCEMENT AND RECOGNITION	6277
<i>Zhaoxu Nian, Jun Du, University of Science and Technology of China, China; Yu Ting Yeung, Renyu Wang, Huawei Noah's Ark Lab, China</i>	
SPE-10.2: A TWO-STEP APPROACH TO LEVERAGE CONTEXTUAL DATA: SPEECH RECOGNITION IN AIR-TRAFFIC COMMUNICATION	6282
<i>Iuliia Nigmatulina, Juan Zuluaga-Gomez, Amrutha Prasad, Seyyed Saeed Sarfjoo, Petr Motlicek, Idiap Research Institute, Switzerland</i>	
SPE-10.3: LEARNING TO ENHANCE OR NOT: NEURAL NETWORK-BASED SWITCHING OF ENHANCED AND OBSERVED SIGNALS FOR OVERLAPPING SPEECH RECOGNITION	6287
<i>Hiroshi Sato, Tsubasa Ochiai, Marc Delcroix, Keisuke Kinoshita, Naoyuki Kamo, Takafumi Moriya, NTT Corporation, Japan</i>	
SPE-10.4: INTERACTIVE FEATURE FUSION FOR END-TO-END NOISE-ROBUST SPEECH RECOGNITION	6292
<i>Yuchen Hu, Nana Hou, Chen Chen, Eng Siong Chng, Nanyang Technological University, Singapore</i>	
SPE-10.5: SPEAKER REINFORCEMENT USING TARGET SOURCE EXTRACTION FOR ROBUST AUTOMATIC SPEECH RECOGNITION	6297
<i>Catalin Zorila, Rama Doddipatla, Toshiba Cambridge Research Laboratory, United Kingdom of Great Britain and Northern Ireland</i>	
SPE-10.6: MITIGATING CLOSED-MODEL ADVERSARIAL EXAMPLES WITH BAYESIAN NEURAL MODELING FOR ENHANCED END-TO-END SPEECH RECOGNITION	6302
<i>Chao-Han Huck Yang, Georgia Institute of Technology, United States of America; Zeeshan Ahmed, Yile Gu, Joseph Szurley, Roger Ren, Linda Liu, Andreas Stolcke, Ivan Bulyko, Amazon, United States of America</i>	
 SPE-11: SPEECH SYNTHESIS: STYLE & EXPRESSIVENESS	
SPE-11.1: REFEREE: TOWARDS REFERENCE-FREE CROSS-SPEAKER STYLE TRANSFER WITH LOW-QUALITY DATA FOR EXPRESSIVE SPEECH SYNTHESIS	6307
<i>Songxiang Liu, Shan Yang, Dan Su, Dong Yu, Tencent, China</i>	
SPE-11.2: PVAE-TTS: ADAPTIVE TEXT-TO-SPEECH VIA PROGRESSIVE STYLE ADAPTATION	6312
<i>Ji-Hyun Lee, Sang-Hoon Lee, Ji-Hoon Kim, Seong-Whan Lee, Korea University, Korea, Republic of</i>	

SPE-11.3: EMOQ-TTS: EMOTION INTENSITY QUANTIZATION FOR FINE-GRAINED CONTROLLABLE EMOTIONAL TEXT-TO-SPEECH	6317
<i>Chae-Bin Im, Sang-Hoon Lee, Seung-Bin Kim, Seong-Whan Lee, Korea University, Korea, Republic of</i>	
SPE-11.4: JOINT AND ADVERSARIAL TRAINING WITH ASR FOR EXPRESSIVE SPEECH SYNTHESIS	6322
<i>Kaili Zhang, Cheng Gong, Wenhuan Lu, Longbiao Wang, Jianguo Wei, Dawei Liu, Tianjin University, China</i>	
SPE-11.5: MSDTRON: A HIGH-CAPABILITY MULTI-SPEAKER SPEECH SYNTHESIS SYSTEM FOR DIVERSE DATA USING CHARACTERISTIC INFORMATION	6327
<i>Qinghua Wu, Quanbo Shen, Jian Luan, Yujun Wang, Xiaomi Company, China</i>	
SPE-11.6: SPEECHSPLIT2.0: UNSUPERVISED SPEECH DISENTANGLEMENT FOR VOICE CONVERSION WITHOUT TUNING AUTOENCODER BOTTLENECKS	6332
<i>Chak Ho Chan, Mark Hasegawa-Johnson, University of Illinois at Urbana-Champaign, United States of America; Kaizhi Qian, Yang Zhang, MIT-IBM Watson AI Lab, United States of America</i>	
SPE-12: EVENT AND RELATION EXTRACTION	
SPE-12.1: DOCUMENT-LEVEL EVENT EXTRACTION VIA HUMAN-LIKE READING PROCESS	6337
<i>Shiyao Cui, Xin Cong, Bowen Yu, Tingwen Liu, Yucheng Wang, Institute of Information Engineering, Chinese Academy of Sciences., China; Jinqiao Shi, Beijing University of Posts and Telecommunications, China</i>	
SPE-12.2: GENERATING DISENTANGLED ARGUMENTS WITH PROMPTS: A SIMPLE EVENT EXTRACTION FRAMEWORK THAT WORKS	6342
<i>Jinghui Si, Beihang University; Alibaba Group, China; Xutan Peng, The University of Sheffield, United Kingdom of Great Britain and Northern Ireland; Chen Li, Jianxin Li, Beihang University, China; Haotian Xu, Alibaba Group, China</i>	
SPE-12.3: MULTI-ROLE EVENT ARGUMENT EXTRACTION AS MACHINE READING COMPREHENSION WITH ARGUMENT MATCH OPTIMIZATION	6347
<i>Jingcong Tao, Youcheng Pan, Baotian Hu, Xiaolong Wang, Harbin Institute of Technology, Shenzhen, China; Xinyu Li, Weihua Peng, Cuiyun Han, Baidu International Technology (Shenzhen), China</i>	
SPE-12.4: BNU: A BALANCE-NORMALIZATION-UNCERTAINTY MODEL FOR INCREMENTAL EVENT DETECTION	6352
<i>Jia Li, Peking University, China; Yunyan Zhang, Yifan Yang, Yefeng Zheng, Tencent, China; Zhicheng An, Tsinghua-Berkeley Shenzhen Institute, Tsinghua University, China</i>	
SPE-12.5: WLINKER: MODELING RELATIONAL TRIPLET EXTRACTION AS WORD LINKING	6357
<i>Yongxiu Xu, Chuan Zhou, Jing Yu, Yue Hu, University of Chinese Academy of Sciences, China; Heyan Huang, Beijing Institute of Technology, China</i>	
SPE-12.6: A KNOWLEDGE/DATA ENHANCED METHOD FOR JOINT EVENT AND TEMPORAL RELATION EXTRACTION	6362
<i>Xiaobin Zhang, Liangjun Zang, Songlin Hu, Institute of Information Engineering, Chinese Academy of Sciences, China; Peng Cheng, Yuqi Wang, State Key Laboratory of Media Convergence Production Technology and Systems, China</i>	
SPE-13: SPEAKER RECOGNITION II: ANTI-SPOOFING	
SPE-13.1: AASIST: AUDIO ANTI-SPOOFING USING INTEGRATED SPECTRO-TEMPORAL GRAPH ATTENTION NETWORKS	6367
<i>Jee-weon Jung, Hee-Soo Heo, Joon Son Chung, Bong-Jin Lee, Naver Corporation, Korea, Republic of; Hemlata Tak, Nicholas Evans, EURECOM, France; Hye-jin Shim, Ha-Jin Yu, University of Seoul, France</i>	

SPE-13.2: ESTIMATING THE CONFIDENCE OF SPEECH SPOOFING COUNTERMEASURE	6372
<i>Xin Wang, Junichi Yamagishi, National Institute of Informatics, Japan</i>	
SPE-13.3: TWO-PATH GMM-RESNET AND GMM-SENET FOR ASV SPOOFING DETECTION	6377
<i>Zhenchun Lei, Hui Yan, Changhong Liu, Minglei Ma, Yingen Yang, Jiangxi Normal University, China</i>	
SPE-13.4: RAWBOOST: A RAW DATA BOOSTING AND AUGMENTATION METHOD APPLIED TO AUTOMATIC SPEAKER VERIFICATION ANTI-SPOOFING	6382
<i>Hemlata Tak, Madhu Kamble, Jose Patino, Massimiliano Todisco, Nicholas Evans, EURECOM, France</i>	
SPE-13.5: EXPLAINING DEEP LEARNING MODELS FOR SPOOFING AND DEEPFAKE DETECTION WITH SHAPLEY ADDITIVE EXPLANATIONS	6387
<i>Wanying Ge, Jose Patino, Massimiliano Todisco, Nicholas Evans, EURECOM, France</i>	
SPE-13.6: MULTI-TASK LEARNING IMPROVES SYNTHETIC SPEECH DETECTION	6392
<i>Yichuan Mo, Shilin Wang, Shanghai Jiao Tong University, China</i>	
 SPE-14: MULTI-LINGUAL ASR	
SPE-14.1: MASSIVELY MULTILINGUAL ASR: A LIFELONG LEARNING SOLUTION	6397
<i>Bo Li, Ruoming Pang, Yu Zhang, Tara Sainath, Trevor Strohman, Parisa Haghani, Yun Zhu, Brian Farris, Neeraj Gaur, Manasa Prasad, Google, United States of America</i>	
SPE-14.2: JOINT UNSUPERVISED AND SUPERVISED TRAINING FOR MULTILINGUAL ASR	6402
<i>Junwen Bai, Cornell University, United States of America; Bo Li, Yu Zhang, Ankur Bapna, Nikhil Siddhartha, Khe Chai Sim, Tara Sainath, Google, United States of America</i>	
SPE-14.3: MULTILINGUAL SECOND-PASS RESCORING FOR AUTOMATIC SPEECH RECOGNITION SYSTEMS	6407
<i>Neeraj Gaur, Tongzhou Chen, Ehsan Variani, Parisa Haghani, Bhuvana Ramabhadran, Pedro Moreno, Google, United States of America</i>	
SPE-14.4: JOINT MODELING OF CODE-SWITCHED AND MONOLINGUAL ASR VIA CONDITIONAL FACTORIZATION	6412
<i>Brian Yan, Siddharth Dalmia, Dan Berrebbi, Shinji Watanabe, Carnegie Mellon University, United States of America; Chunlei Zhang, Meng Yu, Shi-Xiong Zhang, Chao Weng, Dong Yu, Tencent AI Lab, United States of America</i>	
SPE-14.5: BILINGUAL END-TO-END ASR WITH BYTE-LEVEL SUBWORDS	6417
<i>Liuhui Deng, Roger Hsiao, Arnab Ghoshal, Apple, United States of America</i>	
SPE-14.6: A CONFIGURABLE MULTILINGUAL MODEL IS ALL YOU NEED TO RECOGNIZE ALL LANGUAGES	6422
<i>Long Zhou, Shujie Liu, Microsoft Research Asia, China; Jinyu Li, Eric Sun, Microsoft Speech and Language Group, United States of America</i>	
 SPE-15: EMOTION RECOGNITION: REPRESENTATION LEARNING	
SPE-15.1: DOMAIN-INVARIANT FEATURE LEARNING FOR CROSS CORPUS SPEECH EMOTION RECOGNITION	6427
<i>Yuan Gao, Longbiao Wang, Jiaxing Liu, Jianwu Dang, Tianjin University, China; Shogo Okada, Japan Advanced Institute of Science and Technology, Japan</i>	

SPE-15.2: MULTI-STAGE GRAPH REPRESENTATION LEARNING FOR DIALOGUE-LEVEL SPEECH EMOTION RECOGNITION	6432
<i>Yaodong Song, Jiaying Liu, Longbiao Wang, Ruiguo Yu, Tianjin University, China; Jianwu Dang, Japan Advanced Institute of Science and Technology, Japan</i>	
SPE-15.3: SPEECH EMOTION RECOGNITION WITH GLOBAL-AWARE FUSION ON MULTI-SCALE FEATURE REPRESENTATION	6437
<i>Wenjing Zhu, Xiang Li, DXM Financial, China</i>	
SPE-15.4: REPRESENTATION LEARNING THROUGH CROSS-MODAL CONDITIONAL TEACHER-STUDENT TRAINING FOR SPEECH EMOTION RECOGNITION	6442
<i>Sundararajan Srinivasan, Zhaocheng Huang, Katrin Kirchhoff, Amazon.com, United States of America</i>	
SPE-15.5: NOT ALL FEATURES ARE EQUAL: SELECTION OF ROBUST FEATURES FOR SPEECH EMOTION RECOGNITION IN NOISY ENVIRONMENTS	6447
<i>Seong-Gyun Leem, Carlos Busso, The University of Texas at Dallas, United States of America; Daniel Fulford, Boston University, United States of America; Jukka-Pekka Onnela, Harvard University, United States of America; David Gard, San Francisco State University, United States of America</i>	
SPE-15.6: TOWARDS TRANSFERABLE SPEECH EMOTION REPRESENTATION: ON LOSS FUNCTIONS FOR CROSS-LINGUAL LATENT REPRESENTATIONS	6452
<i>Sneha Das, Line H. Clemmensen, Technical University of Denmark, Denmark; Nicole Nadine Lønfeldt, Copenhagen University Hospital, Capital Region, Denmark; Anne Katrine Pagsberg, Copenhagen University Hospital and Copenhagen University, Denmark</i>	
SPE-16: LANGUAGE DISORDERS: DETECTION I	
SPE-16.1: DEMENTIA DETECTION BY FUSING SPEECH AND EYE-TRACKING REPRESENTATION	6457
<i>Zhengyan Sheng, Zhiqiang Guo, University of Science and Technology of China, China; Xin Li, iFlytek, China; Yunxia Li, Shanghai Tongji Hospital, China; Zhenhua Ling, National Engineering Laboratory for Speech and Language Information Processing, University of Science and Technology of China, Hefei, P.R.China, China</i>	
SPE-16.2: TOWARDS INTERPRETABILITY OF SPEECH PAUSE IN DEMENTIA DETECTION USING ADVERSARIAL LEARNING	6462
<i>Youxiang Zhu, Bang Tran, Xiaohui Liang, University of Massachusetts Boston, United States of America; John Batsis, University of North Carolina at Chapel Hill, United States of America; Robert Roth, Geisel School of Medicine at Dartmouth, United States of America</i>	
SPE-16.3: USING SPECTRAL SEQUENCE-TO-SEQUENCE AUTOENCODERS TO ASSESS MILD COGNITIVE IMPAIRMENT	6467
<i>Mercedes Vetráb, José Vicente Egas-López, Réka Balogh, Nóra Imre, László Tóth, Magdolna Pákáski, János Kálmán, University of Szeged, Hungary; Ildikó Hoffmann, Research Institute for Linguistics, Hungary; Gábor Gosztolya, ELRN-SZTE Research Group on Artificial Intelligence, Hungary</i>	
SPE-16.4: EXPLORING DEMENTIA DETECTION FROM SPEECH: CROSS CORPUS ANALYSIS	6472
<i>Ayimmisagul Ablimit, Tanja Schultz, University of Bremen, Germany; Catarina Botelho, Alberto Abad, Isabel Trancoso, INESC-ID/Instituto Superior Técnico, Portugal</i>	
SPE-16.5: EXPERIMENTAL INVESTIGATION ON STFT PHASE REPRESENTATIONS FOR DEEP LEARNING-BASED DYSARTHIC SPEECH DETECTION	6477
<i>Parvaneh Janbakhshi, Ina Kodrasi, Idiap Research Institute, Switzerland</i>	
SPE-16.6: DYSFLUENCY CLASSIFICATION IN STUTTERED SPEECH USING DEEP LEARNING FOR REAL-TIME APPLICATIONS	6482
<i>Melanie Jouaiti, Kerstin Dautenhahn, University of Waterloo, Canada</i>	

SPE-17: SPEECH ENHANCEMENT: MULTI-CHANNEL

SPE-17.1: EMBEDDING AND BEAMFORMING: ALL-NEURAL CAUSAL BEAMFORMER 6487 FOR MULTICHANNEL SPEECH ENHANCEMENT

Andong Li, Wenzhe Liu, Chengshi Zheng, Xiaodong Li, Institute of Acoustics, Chinese Academy of Sciences, China

SPE-17.2: IMPROVING DUAL-MICROPHONE SPEECH ENHANCEMENT BY 6492 LEARNING CROSS-CHANNEL FEATURES WITH MULTI-HEAD ATTENTION

Xinmeng Xu, Rongzhi Gu, Yuexian Zou, Peking University, China

SPE-17.3: TPARN: TRIPLE-PATH ATTENTIVE RECURRENT NETWORK FOR 6497 TIME-DOMAIN MULTICHANNEL SPEECH ENHANCEMENT

Ashutosh Pandey, DeLiang Wang, The Ohio State University, United States of America; Buye Xu, Anurag Kumar, Jacob Donley, Paul Calamia, Facebook Reality Labs Research, United States of America

SPE-17.4: MULTICHANNEL SPEECH ENHANCEMENT WITHOUT BEAMFORMING 6502

Ashutosh Pandey, DeLiang Wang, The Ohio State University, United States of America; Buye Xu, Anurag Kumar, Jacob Donley, Paul Calamia, Facebook Reality Labs Research, United States of America

SPE-17.5: LEARNING FILTERBANKS FOR END-TO-END ACOUSTIC BEAMFORMING 6507

Samuele Cornell, Stefano Squartini, Università Politecnica delle Marche, Italy; Manuel Pariente, Université de Lorraine, CNRS, Inria, LORIA, France; François Grondin, Université de Sherbrooke, Canada

SPE-17.6: SPATIAL-TEMPORAL GRAPH CONVOLUTION NETWORK FOR 6512 MULTICHANNEL SPEECH ENHANCEMENT

Minghui Hao, Jingjing Yu, Luyao Zhang, Beijing Jiaotong University, China

SPE-18: SPEECH RECOGNITION: GENERAL TOPICS I

SPE-18.1: LATTICE RESCORING BASED ON LARGE ENSEMBLE OF 6517 COMPLEMENTARY NEURAL LANGUAGE MODELS

Atsunori Ogawa, Naohiro Tawara, Marc Delcroix, Shoko Araki, NTT Corporation, Japan

SPE-18.2: CONTINUAL LEARNING USING LATTICE-FREE MMI FOR SPEECH 6522 RECOGNITION

Hossein Hadian, Arseniy Gorin, Behavox, Canada

SPE-18.3: NON-AUTOREGRESSIVE TRANSFORMER WITH UNIFIED 6527 BIDIRECTIONAL DECODER FOR AUTOMATIC SPEECH RECOGNITION

Chuan-Fei Zhang, Yan Liu, Tian-Hao Zhang, Song-Lu Chen, Xu-Cheng Yin, University of Science and Technology Beijing, China; Feng Chen, EEasy Technology Company Ltd., China

SPE-18.4: MODEL-BASED APPROACH FOR MEASURING THE FAIRNESS IN ASR 6532

Zhe Liu, Irina-Elena Veliche, Fuchun Peng, Meta Platforms, Inc, United States of America

SPE-18.5: IMPROVING CONFIDENCE ESTIMATION ON OUT-OF-DOMAIN DATA FOR 6537 END-TO-END SPEECH RECOGNITION

Qiujia Li, Phil Woodland, University of Cambridge, United Kingdom of Great Britain and Northern Ireland; Yu Zhang, David Qiu, Yanzhang He, Liangliang Cao, Google LLC, United States of America

SPE-18.6: PARALLEL COMPOSITION OF WEIGHTED FINITE-STATE TRANSDUCERS 6542

Shubho Sengupta, Vineel Pratap, Facebook, United States of America; Awni Hannun, Zoom, United States of America

SPE-19: VOICE CONVERSION: REPRESENTATION

SPE-19.1: DGC-VECTOR: A NEW SPEAKER EMBEDDING FOR ZERO-SHOT VOICE 6547 CONVERSION

Ruitong Xiao, South China University of Technology, China; Haitong Zhang, Yue Lin, Netease Games, China

SPE-19.2: S3PRL-VC: OPEN-SOURCE VOICE CONVERSION FRAMEWORK WITH SELF-SUPERVISED SPEECH REPRESENTATIONS	6552
<i>Wen-Chin Huang, Tomoki Hayashi, Tomoki Toda, Nagoya University, Japan; Shu-wen Yang, Hung-yi Lee, National Taiwan University, Taiwan; Shinji Watanabe, Carnegie Mellon University, United States of America</i>	
SPE-19.3: TRAINING ROBUST ZERO-SHOT VOICE CONVERSION MODELS WITH SELF-SUPERVISED FEATURES	6557
<i>Trung Dang, Peter Chin, Boston University, United States of America; Dung Tran, Kazuhito Koishida, Microsoft Corp., United States of America</i>	
SPE-19.4: A COMPARISON OF DISCRETE AND SOFT SPEECH UNITS FOR IMPROVED VOICE CONVERSION	6562
<i>Benjamin van Niekerk, Matthew Baas, Herman Kamper, Stellenbosch University, South Africa; Marc-André Carbonneau, Julian Zaïdi, Hugo Seuté, Ubisoft, Canada</i>	
SPE-19.5: SIG-VC: A SPEAKER INFORMATION GUIDED ZERO-SHOT VOICE CONVERSION SYSTEM FOR BOTH HUMAN BEINGS AND MACHINES	6567
<i>Haozhe Zhang, Zexin Cai, Xiaoyi Qin, Ming Li, Duke Kunshan University, China</i>	
SPE-19.6: ROBUST DISENTANGLED VARIATIONAL SPEECH REPRESENTATION LEARNING FOR ZERO-SHOT VOICE CONVERSION	6572
<i>Jiachen Lian, UC Berkeley, United States of America; Chunlei Zhang, Dong Yu, Tencent AI Lab, United States of America</i>	
 SPE-20: DIALOG SYSTEM I: RESPONSE RETRIEVAL/GENERATION	
SPE-20.1: IMPROVING DIALOGUE GENERATION VIA PROACTIVELY QUERYING GROUNDED KNOWLEDGE	6577
<i>Xiangyu Zhao, Longbiao Wang, Tianjin University, China; Jianwu Dang, Japan Advanced Institute of Science and Technology, China; , ,</i>	
SPE-20.2: A NON-HIERARCHICAL ATTENTION NETWORK WITH MODALITY DROPOUT FOR TEXTUAL RESPONSE GENERATION IN MULTIMODAL DIALOGUE SYSTEMS	6582
<i>Rongyi Sun, Borun Chen, Yinghui Li, Hai-Tao Zheng, Shenzhen International Graduate School, Tsinghua University, China; Qingyu Zhou, Yunbo Cao, Tencent Cloud Xiaowei, China; , ,</i>	
SPE-20.3: JOINT LEARNING FOR ADDRESSEE SELECTION AND RESPONSE GENERATION IN MULTI-PARTY CONVERSATION	6587
<i>Qi Song, Sheng Li, Ping Wei, Ge Luo, Xinpeng Zhang, Zhenxing Qian, Fudan University, China</i>	
SPE-20.4: RETRIEVAL ENHANCED SEGMENT GENERATION NEURAL NETWORK FOR TASK-ORIENTED DIALOGUE SYSTEMS	6592
<i>Miaoxin Chen, Rongyi Sun, Kai Ouyang, Hai-Tao Zheng, Shenzhen International Graduate School, Tsinghua University, China, China; Zibo Lin, WeChat Search Application Department, Tencent, China, China; Rui Xie, Wei Wu, Meituan, Shanghai, China, China</i>	
SPE-20.5: A MULTI DOMAIN KNOWLEDGE ENHANCED MATCHING NETWORK FOR RESPONSE SELECTION IN RETRIEVAL-BASED DIALOGUE SYSTEMS	6597
<i>Xiuyi Chen, Feilong Chen, Shuang Xu, Bo Xu, Institute of Automation, Chinese Academy of Science, China</i>	
SPE-20.6: RETRIEVAL BIAS AWARE ENSEMBLE MODEL FOR CONDITIONAL SENTENCE GENERATION	6602
<i>Yiping Song, Zheng Xie, Jianping Li, National University of Defense Technology, China; Luchen Liu, Ming Zhang, Peking University, China; Zhiliang Tian, The Hong Kong University of Science and Technology, China</i>	

SPE-21: SPEAKER RECOGNITION III: ADVERSARIAL, SPOOFING AND ROBUSTNESS

SPE-21.2: EFFECTIVE AND INCONSPICUOUS OVER-THE-AIR ADVERSARIAL 6607 EXAMPLES WITH ADAPTIVE FILTERING

Patrick O'Reilly, Aravindan Vijayaraghavan, Bryan Pardo, Northwestern University, United States of America; Pranjal Awasthi, Google Research, United States of America

SPE-21.3: LRPD: LARGE REPLAY PARALLEL DATASET 6612

Ivan Yakovlev, Mikhail Melnikov, Nikita Bukhal, Rostislav Makarov, Alexander Alenin, Nikita Torgashov, Anton Okhotnikov, ID R&D Inc., Russian Federation

SPE-21.4: ROBUST SELF-SUPERVISED SPEAKER REPRESENTATION LEARNING VIA 6617 INSTANCE MIX REGULARIZATION

Woohyun Kang, Jahangir Alam, Abderrahim Fathan, Computer Research Institute of Montreal, Canada

SPE-21.5: GRAPH CONVOLUTIONAL NETWORK BASED SEMI-SUPERVISED 6622 LEARNING ON MULTI-SPEAKER MEETING DATA

Fuchuan Tong, Lin Li, Qingyang Hong, Xiamen University, China; Siqi Zheng, Hongbin Suo, Alibaba Group, China; Min Zhang, Zhejiang University, China; Yafeng Chen, University of Science and Technology of China, China

SPE-22: ADAPTATION AND PERSONALIZATION FOR SPEECH MODELS

SPE-22.1: LARGE-SCALE ASR DOMAIN ADAPTATION USING SELF- AND 6627 SEMI-SUPERVISED LEARNING

Dongseong Hwang, Ananya Misra, Zhouyuan Huo, Nikhil Siddhartha, Shefali Garg, David Qiu, Khe Chai Sim, Trevor Strohman, Françoise Beaufays, Yanzhang He, Google, United States of America

SPE-22.2: FAST CONTEXTUAL ADAPTATION WITH NEURAL ASSOCIATIVE MEMORY 6632 FOR ON-DEVICE PERSONALIZED SPEECH RECOGNITION

Tsendsuren Munkhdalai, Khe Chai Sim, Angad Chandorkar, Fan Gao, Mason Chua, Trevor Strohman, Françoise Beaufays, Google, United States of America

SPE-22.3: PERSONALIZED AUTOMATIC SPEECH RECOGNITION TRAINED ON 6637 SMALL DISORDERED SPEECH DATASETS

Jimmy Tobin, Katrin Tomanek, Google, United States of America

SPE-22.4: SPELL MY NAME: KEYWORD BOOSTED SPEECH RECOGNITION..... 6642

Namkyu Jung, Geonmin Kim, NAVER Corporation, Korea, Republic of; Joon Son Chung, Korea Advanced Institute of Science and Technology, Korea, Republic of

SPE-22.5: MAGIC DUST FOR CROSS-LINGUAL ADAPTATION OF MONOLINGUAL 6647 WAV2VEC-2.0

Sameer Khurana, James Glass, Massachusetts Institute of Technology, United States of America; Antoine Laurent, Le Mans University, France

SPE-22.6: TEXT ADAPTIVE DETECTION FOR CUSTOMIZABLE KEYWORD 6652 SPOTTING

Yu Xi, Wangyou Zhang, Baochen Yang, Kai Yu, Shanghai Jiao Tong University, China; Tian Tan, AISpeech Ltd, China

SPE-23: VOICE CONVERSION: SINGING VOICE AND OTHERS

SPE-23.1: IMPROVING ADVERSARIAL WAVEFORM GENERATION BASED SINGING 6657 VOICE CONVERSION WITH HARMONIC SIGNALS

Haohan Guo, Chinese University of Hong Kong, Hong Kong; Zhiping Zhou, Fanbo Meng, Kai Liu, Sogou, China

SPE-23.2: K-CONVERTER: AN UNSUPERVISED SINGING VOICE CONVERSION 6662 SYSTEM

Ying Zhang, Peng Yang, Jinba Xiao, Ye Bai, Hao Che, Xiaorui Wang, kwai, China

SPE-23.3: HIFI-SVC: FAST HIGH FIDELITY CROSS-DOMAIN SINGING VOICE CONVERSION	6667
<i>Yong Zhou, Xiangju Lu, iQIYI Inc., China</i>	
SPE-23.4: TOWARDS IDENTITY PRESERVING NORMAL TO DYSARTHIC VOICE CONVERSION	6672
<i>Wen-Chin Huang, Lester Phillip Violeta, Tomoki Toda, Nagoya University, Japan; Bence Mark Halpern, Odette Scharenborg, Delft University of Technology, Netherlands</i>	
SPE-23.5: SPEAKER IDENTITY PRESERVATION IN DYSARTHIC SPEECH RECONSTRUCTION BY ADVERSARIAL SPEAKER ADAPTATION	6677
<i>Disong Wang, Songxiang Liu, Xixin Wu, Hui Lu, Xunying Liu, Helen Meng, The Chinese University of Hong Kong, Hong Kong; Lifa Sun, SpeechX Limited, China</i>	
SPE-23.6: CONTROLLABLE SPEECH REPRESENTATION LEARNING VIA VOICE CONVERSION AND AIC LOSS	6682
<i>Yunyun Wang, Jiaqi Su, Adam Finkelstein, Princeton University, United States of America; Zeyu Jin, Adobe Research, United States of America</i>	
SPE-24: DIALOG SYSTEM II: GENERAL TOPICS	
SPE-24.1: VU-BERT: A UNIFIED FRAMEWORK FOR VISUAL DIALOG	6687
<i>Tong Ye, Ping An Technology (Shenzhen) Co., Ltd. ;University of Science and Technology of China, China; Shijing Si, Jianzong Wang, Ning Cheng, Jing Xiao, Ping An Technology (Shenzhen) Co., Ltd., China; Rui Wang, Duke University, China</i>	
SPE-24.2: INTEGRATING PRETRAINED LANGUAGE MODEL FOR DIALOGUE POLICY EVALUATION	6692
<i>Hongru Wang, Huimin Wang, Zezhong Wang, Kam-Fai Wong, The Chinese University of Hong Kong, China</i>	
SPE-24.3: CACHE: MODELING CONTRIBUTION-AWARE CONTEXT HIERARCHICALLY FOR LONG-RANGE DIALOGUE STATE TRACKING	6697
<i>Jianshu Qi, Yuke Si, Longbiao Wang, Tianjin University, China; Jianwu Dang, Japan Advanced Institute of Science and Technology, Japan</i>	
SPE-24.4: ROBUST UNSTRUCTURED KNOWLEDGE ACCESS IN CONVERSATIONAL DIALOGUE WITH ASR ERRORS	6702
<i>Yik-Cheung Tam, Jiacheng Xu, Zecheng Wang, Tinglong Liao, Shuhan Yuan, NYU Shanghai, China; Jiakai Zou, NYU, China</i>	
SPE-24.5: AN EMBARRASSINGLY SIMPLE MODEL FOR DIALOGUE RELATION EXTRACTION	6707
<i>Fuzhao Xue, Aixin Sun, Hao Zhang, Jinjie Ni, Eng-Siong Chng, Nanyang Technological University, Singapore</i>	
SPE-24.6: A GAUSSIAN MIXTURE MODEL FOR DIALOGUE GENERATION WITH DYNAMIC PARAMETER SHARING STRATEGY	6712
<i>Qingqing Zhu, Pengfei Wu, Zhouxing Tan, Jiaxin Duan, Fengyu Lu, Junfei Liu, Peking University, China</i>	
SPE-25: SPEAKER RECOGNITION IV: ATTENTION MECHANISM	
SPE-25.1: ATTENTION BACK-END FOR AUTOMATIC SPEAKER VERIFICATION WITH MULTIPLE ENROLLMENT UTTERANCES	6717
<i>Chang Zeng, Junichi Yamagishi, National Institute of Informatics, SOKENDAI, Japan; Xin Wang, Erica Cooper, Xiaoxiao Miao, National Institute of Informatics, Japan</i>	

SPE-25.2: SIMPLE ATTENTION MODULE BASED SPEAKER VERIFICATION WITH ITERATIVE NOISY LABEL DETECTION	6722
<i>Xiaoyi Qin, Wuhan University, China; Na Li, Chao Weng, Dan Su, Tencent AI Lab, China; Ming Li, Duke Kunshan University, China</i>	
SPE-25.3: LOCAL INFORMATION MODELING WITH SELF-ATTENTION FOR SPEAKER VERIFICATION	6727
<i>Bing Han, Zhengyang Chen, Yanmin Qian, Shanghai Jiao Tong University, China</i>	
SPE-25.4: MULTI-VIEW SELF-ATTENTION BASED TRANSFORMER FOR SPEAKER RECOGNITION	6732
<i>Rui Wang, Zhihua Wei, Tongji University, China; Junyi Ao, Tom Ko, Yu Zhang, Southern University of Science and Technology, China; Long Zhou, Shujie Liu, Microsoft Research Asia, China; Qing Li, The Hong Kong Polytechnic University, Hong Kong</i>	
SPE-25.5: MULTI-QUERY MULTI-HEAD ATTENTION POOLING AND INTER-TOPK PENALTY FOR SPEAKER VERIFICATION	6737
<i>Miao Zhao, Yufeng Ma, Yu Zheng, Min Liu, Minqiang Xu, SpeakIn Technologies Co. Ltd., China; Yiwei Ding, Fudan University, China</i>	
SPE-25.6: TEMPORAL DYNAMIC CONVOLUTIONAL NEURAL NETWORK FOR TEXT-INDEPENDENT SPEAKER VERIFICATION AND PHONEMIC ANALYSIS	6742
<i>Seong-Hu Kim, Hyeonuk Nam, Yong-Hwa Park, Korea Advanced Institute of Science and Technology, Korea, Republic of</i>	
 SPE-26: SPEECH RECOGNITION: GENERAL TOPICS II	
SPE-26.1: EXPLOITING CROSS DOMAIN ACOUSTIC-TO-ARTICULATORY INVERTED FEATURES FOR DISORDERED SPEECH RECOGNITION	6747
<i>Shujie Hu, Shansong Liu, Xurong Xie, Mengzhe Geng, Tianzi Wang, Shoukang Hu, Mingyu Cui, Xunying Liu, Helen Meng, The Chinese University of Hong Kong, Hong Kong</i>	
SPE-26.2: CONVERSATIONAL SPEECH RECOGNITION BY LEARNING CONVERSATION-LEVEL CHARACTERISTICS	6752
<i>Kun Wei, Lei Xie, Audio, Speech and Language Processing Group, School of Computer Science, Northwestern Polytechnical University, China; Yike Zhang, Sining Sun, Long Ma, Tencent, xiaowei, China</i>	
SPE-26.3: EXPLORING MACHINE SPEECH CHAIN FOR DOMAIN ADAPTATION	6757
<i>Fengpeng Yue, Tom Ko, Yu Zhang, Southern University of Science and Technology, China; Yan Deng, Lei He, Microsoft China, China</i>	
SPE-26.4: A LIKELIHOOD RATIO BASED DOMAIN ADAPTATION METHOD FOR E2E MODELS	6762
<i>Chhavi Choudhury, Ankur Gandhe, Xiaohan Ding, Ivan Bulyko, Amazon, United States of America</i>	
SPE-26.5: RETRIEVING SPEAKER INFORMATION FROM PERSONALIZED ACOUSTIC MODELS FOR SPEECH RECOGNITION	6767
<i>Salima Mdhaffar, Jean-François Bonastre, Natalia Tomashenko, Yannick Estève, LIA - University of Avignon, France; Marc Tommasi, Inria - University of Lille, France</i>	
SPE-26.6: NON-AUTOREGRESSIVE END-TO-END AUTOMATIC SPEECH RECOGNITION INCORPORATING DOWNSTREAM NATURAL LANGUAGE PROCESSING	6772
<i>Motoi Omachi, Yuya Fujita, Yahoo Japan Corporation, Japan; Shinji Watanabe, Carnegie Mellon University, Japan; Tianzi Wang, Johns Hopkins University, Japan</i>	

SPE-27: VOICE CONVERSION I

SPE-27.1: TOWARD DEGRADATION-ROBUST VOICE CONVERSION 6777
Chien-yu Huang, Kai-Wei Chang, Hung-yi Lee, National Taiwan University, Taiwan

SPE-27.2: TEXT-FREE NON-PARALLEL MANY-TO-MANY VOICE CONVERSION 6782
USING NORMALISING FLOWS
Thomas Merritt, Abdelhamid Ezzer, Piotr Bilinski, Kamil Pokora, Roberto Barra-Chicote, Daniel Korzekwa, Amazon, United Kingdom of Great Britain and Northern Ireland; Magdalena Proszewska, Jagiellonian University, Poland

SPE-27.3: DIRECT NOISY SPEECH MODELING FOR NOISY-TO-NOISY VOICE 6787
CONVERSION
Chao Xie, Yi-Chiao Wu, Patrick Lumban Tobing, Wen-Chin Huang, Tomoki Toda, Nagoya University, Japan

SPE-27.4: ONE-SHOT VOICE CONVERSION FOR STYLE TRANSFER BASED ON 6792
SPEAKER ADAPTATION
Zhichao Wang, Qicong Xie, Tao Li, Hongqiang Du, Lei Xie, Northwestern Polytechnical University, China; Pengcheng Zhu, Mengxiao Bi, Fuxi AI Lab, NetEase Inc., China

SPE-27.5: CROSS-SPEAKER STYLE TRANSFER FOR TEXT-TO-SPEECH USING DATA 6797
AUGMENTATION
Manuel Sam Ribeiro, Julian Roth, Giulia Comini, Goeric Huybrechts, Adam Gabryś, Jaime Lorenzo-Trueba, Amazon, Poland

SPE-27.6: AN INVESTIGATION OF STREAMING NON-AUTOREGRESSIVE 6802
SEQUENCE-TO-SEQUENCE VOICE CONVERSION
Tomoki Hayashi, Kazuhiro Kobayashi, Tomoki Toda, Nagoya University, Japan

SPE-28: PRONUNCIATION ASSESSMENT

SPE-28.1: A UNIVERSAL ORDINAL REGRESSION FOR ASSESSING 6807
PHONEME-LEVEL PRONUNCIATION
Shaoguang Mao, Frank Soong, Yan Xia, Jonathan Tien, Microsoft Research Asia, China

SPE-28.2: A TRANSFER LEARNING APPROACH FOR PRONUNCIATION SCORING..... 6812
Marcelo Sancinetti, Jazmín Vidal, Cyntia Bonomi, Universidad de Buenos Aires, Argentina; Luciana Ferrer, Consejo Nacional de Investigaciones y Técnicas, Argentina

SPE-28.3: EXPLORING NON-AUTOREGRESSIVE END-TO-END NEURAL MODELING 6817
FOR ENGLISH MISPRONUNCIATION DETECTION AND DIAGNOSIS
Hsin-Wei Wang, Bi-Cheng Yan, Berlin Chen, National Taiwan Normal University, Taiwan; Hsuan-Sheng Chiu, Chunghwa Telecom Laboratories, Taiwan; Yung-Chang Hsu, EZAI, Taiwan

SPE-28.4: PHONEME MISPRONUNCIATION DETECTION BY JOINTLY LEARNING 6822
TO ALIGN
Binghuai Lin, Liyuan Wang, Tencent Technology Co., Ltd, China

SPE-28.5: AN APPROACH TO MISPRONUNCIATION DETECTION AND DIAGNOSIS 6827
WITH ACOUSTIC, PHONETIC AND LINGUISTIC (APL) EMBEDDINGS
Wenxuan Ye, Zhiyong Wu, Tsinghua University, China; Shaoguang Mao, Frank Soong, Wenshan Wu, Yan Xia, Jonathan Tien, Microsoft Research Asia, China

SPE-28.6: MASKED ACOUSTIC UNIT FOR MISPRONUNCIATION DETECTION AND 6832
CORRECTION
Zhan Zhang, Yuehai Wang, Jianyi Yang, Zhejiang University, China

SPE-29: SPEECH ENHANCEMENT AND SEPARATION

SPE-29.1: INVESTIGATING SELF-SUPERVISED LEARNING FOR SPEECH 6837 ENHANCEMENT AND SEPARATION

Zili Huang, Paola Garcia, Sanjeev Khudanpur, Johns Hopkins University, United States of America; Shinji Watanabe, Carnegie Mellon University, United States of America; Shu-wen Yang, National Taiwan University, United States of America

SPE-29.2: TFPSNET: TIME-FREQUENCY DOMAIN PATH SCANNING NETWORK FOR 6842 SPEECH SEPARATION

Lei Yang, Wei Liu, Weiqin Wang, Samsung Research China – Beijing (SRC-B), China

SPE-29.3: EFFICIENT MONAURAL SPEECH SEPARATION WITH MULTISCALE 6847 TIME-DELAY SAMPLING

Shuangqing Qian, Lijian Gao, Hongjie Jia, Qirong Mao, Jiangsu University, China

SPE-29.4: TOWARD MMWAVE-BASED SOUND ENHANCEMENT AND SEPARATION 6852

Muhammed Zahid Ozturk, K. J. Ray Liu, University of Maryland, College Park, United States of America; Chenshu Wu, The University of Hong Kong, Hong Kong; Beibei Wang, Origin Wireless Inc., United States of America

SPE-29.5: DPT-FSNET: DUAL-PATH TRANSFORMER BASED FULL-BAND AND 6857 SUB-BAND FUSION NETWORK FOR SPEECH ENHANCEMENT

Feng Dang, Hangting Chen, Pengyuan Zhang, University of Chinese Academy of Sciences, China

SPE-29.6: REAL-M: TOWARDS SPEECH SEPARATION ON REAL MIXTURES 6862

Cem Subakan, Universite de Sherbrooke / Mila, Canada; Mirco Ravanelli, Concordia University / Mila, Canada; Samuele Cornell, Universita Politecnica delle Marche, Italy; Francois Grondin, Universite de Sherbrooke, Canada

SPE-30: LANGUAGE IDENTIFICATION AND LOW RESOURCE SPEECH RECOGNITION

SPE-30.1: SPOKEN LANGUAGE RECOGNITION WITH CLUSTER-BASED MODELING 6867

Stanisław Kacprzak, Magdalena Rybicka, Konrad Kowalczyk, Akademia Górniczo-Hutnicza im. Stanisława Staszica w Krakowie, Poland

SPE-30.2: PHONOTACTIC LANGUAGE RECOGNITION USING A UNIVERSAL 6872 PHONEME RECOGNIZER AND A TRANSFORMER ARCHITECTURE

David Romero, Christian Salamea, Universidad Politecnica Salesiana, Ecuador; Luis Fernando D'Haro, Marcos Estecha-Garitagotia, Universidad Politécnica de Madrid, Spain

SPE-30.3: IMPROVED LANGUAGE IDENTIFICATION THROUGH CROSS-LINGUAL 6877 SELF-SUPERVISED LEARNING

Andros Tjandra, Diptanu Gon Choudhury, Frank Zhang, Kritika Singh, Alexis Conneau, Alexei Baevski, Assaf Sela, Yatharth Saraf, Michael Auli, Facebook AI, United States of America

SPE-30.4: LANGUAGE ADAPTIVE CROSS-LINGUAL SPEECH REPRESENTATION 6882 LEARNING WITH SPARSE SHARING SUB-NETWORKS

Yizhou Lu, Mingkun Huang, Xinghua Qu, Pengfei Wei, Zejun Ma, ByteDance AI Lab, China

SPE-30.5: INVESTIGATION OF ROBUSTNESS OF HUBERT FEATURES FROM 6887 DIFFERENT LAYERS TO DOMAIN, ACCENT AND LANGUAGE VARIATIONS

Pratik Kumar, Vrunda N. Sukhadia, Srinivasan Umesh, Indian Institute of Technology Madras, India

SPE-30.6: COMBINING UNSUPERVISED AND TEXT AUGMENTED 6892 SEMI-SUPERVISED LEARNING FOR LOW RESOURCED AUTOREGRESSIVE SPEECH RECOGNITION

Chak-Fai Li, Francis Keith, William Hartmann, Matthew Snover, Raytheon BBN Technologies, United States of America

SPE-31: EMOTION RECOGNITION: NEURAL ARCHITECTURE

SPE-31.1: KEY-SPARSE TRANSFORMER FOR MULTIMODAL SPEECH EMOTION RECOGNITION 6897

Weidong Chen, Xiaofeng Xing, Xiangmin Xu, Jichen Yang, South China University of Technology, China; Jianxin Pang, UBTECH Robotics Corp, China

SPE-31.2: NEURAL ARCHITECTURE SEARCH FOR SPEECH EMOTION RECOGNITION 6902

Xixin Wu, Shoukang Hu, Zhiyong Wu, Xunying Liu, Helen Meng, The Chinese University of Hong Kong, Hong Kong

SPE-31.3: MULTI-LINGUAL MULTI-TASK SPEECH EMOTION RECOGNITION USING WAV2VEC 2.0 6907

Mayank Sharma, Amazon, India

SPE-31.4: LIGHT-SERNET: A LIGHTWEIGHT FULLY CONVOLUTIONAL NEURAL NETWORK FOR SPEECH EMOTION RECOGNITION 6912

Arya Aftab, Shahrokh Ghaemmaghami, Sharif University of Technology, Iran (Islamic Republic of); Alireza Morsali, Benoit Champagne, McGill University, Canada

SPE-31.5: MULTIMODAL TRANSFORMER WITH LEARNABLE FRONTEND AND SELF ATTENTION FOR EMOTION RECOGNITION 6917

Soumya Dutta, Sriram Ganapathy, LEAP lab, Indian Institute of Science, Bangalore, India., India

SPE-31.6: SPEECH EMOTION RECOGNITION USING SELF-SUPERVISED FEATURES 6922

Edmilson Moraes, Ron Hoory, Weizhong Zhu, Itai Gat, Matheus Damasceno, Hagai Aronowitz, IBM Research, Brazil

SPE-32: LANGUAGE DISORDERS: DETECTION II

SPE-32.1: USING ACOUSTIC DEEP NEURAL NETWORK EMBEDDINGS TO DETECT MULTIPLE SCLEROSIS FROM SPEECH 6927

Gábor Gosztolya, ELKH-SZTE Research Group on Artificial Intelligence, Hungary; László Tóth, University of Szeged, Hungary; Veronika Svindt, Ildikó Hoffmann, Research Center for Linguistics, Hungary; Judit Bóna, Eötvös Loránd University, Hungary

SPE-32.2: REPETITION ASSESSMENT FOR SPEECH AND LANGUAGE DISORDERS: A STUDY OF THE LOGOPENIC VARIANT OF PRIMARY PROGRESSIVE APHASIA 6932

R'mani Haulcy, James Glass, Massachusetts Institute of Technology, United States of America; Katerina Placek, Brian Tracey, Takeda Pharmaceutical Company, United States of America; Adam Vogel, University of Melbourne/Redenlab Inc, Australia

SPE-32.3: SPEECH TASKS RELEVANT TO SLEEPINESS DETERMINED WITH DEEP TRANSFER LEARNING 6937

Bang Tran, Youxiang Zhu, Xiaohui Liang, University of Massachusetts Boston, United States of America; James Schwoebel, Lindsay Warrenburg, Sonde Health Inc., United States of America

SPE-33: SPEECH ENHANCEMENT: TRAINING SCHEMES AND LOSSES

SPE-33.1: PHASE CONTINUITY: LEARNING DERIVATIVES OF PHASE SPECTRUM FOR SPEECH ENHANCEMENT 6942

Doyeon Kim, Hyewon Han, Hong-Goo Kang, Yonsei University, Korea, Republic of; Hyeon-Kyeong Shin, Soo-Whan Chung, Naver Corporation, Korea, Republic of

SPE-33.2: CONTINUAL SELF-TRAINING WITH BOOTSTRAPPED REMIXING FOR SPEECH ENHANCEMENT	6947
<i>Efthymios Tzinis, University of Illinois at Urbana-Champaign, United States of America; Yossi Adi, Meta AI Research, Israel; Vamsi Ithapu, Buye Xu, Anurag Kumar, Meta Reality Labs Research, United States of America</i>	
SPE-33.3: ALLEVIATING THE LOSS-METRIC MISMATCH IN SUPERVISED SINGLE-CHANNEL SPEECH ENHANCEMENT	6952
<i>Yang Yang, Hui Zhang, Xueliang Zhang, Huaiwen Zhang, Inner Mongolia University, China</i>	
SPE-33.4: A PRIORI SNR ESTIMATION FOR SPEECH ENHANCEMENT BASED ON PESQ-INDUCED REINFORCEMENT LEARNING	6957
<i>Tong Lei, Haoxin Ruan, Kai Chen, Jing Lu, Key Laboratory of Modern Acoustics, Nanjing University; NJU-Horizon Intelligent Audio Lab, Horizon Robotics; Nanjing Institute of Advanced Artificial Intelligence., China</i>	
SPE-33.5: A TRAINING FRAMEWORK FOR STEREO-AWARE SPEECH ENHANCEMENT USING DEEP NEURAL NETWORKS	6962
<i>Bahareh Tolooshams, Harvard University, United States of America; Kazuhito Koishida, Microsoft Corporation, United States of America</i>	
SPE-33.6: JOINT MAGNITUDE ESTIMATION AND PHASE RECOVERY USING CYCLE-IN-CYCLE GAN FOR NON-PARALLEL SPEECH ENHANCEMENT	6967
<i>Guochen Yu, Communication University of China/Institute of Acoustics, Chinese Academy of Sciences, China; Andong Li, Chengshi Zheng, Institute of Acoustics, Chinese Academy of Sciences, China; Yutian Wang, Hui Wang, Communication University of China, China; Yinuo Guo, Bytedance, China</i>	
SPE-34: SPEECH RECOGNITION: GENERAL TOPICS III	
SPE-34.1: PRIVACY ATTACKS FOR AUTOMATIC SPEECH RECOGNITION ACOUSTIC MODELS IN A FEDERATED LEARNING FRAMEWORK	6972
<i>Natalia Tomashenko, Salima Mdhaffar, Yannick Estève, Jean-Francois Bonastre, University of Avignon, France; Marc Tommasi, University of Lille, France</i>	
SPE-34.2: VADOI: VOICE-ACTIVITY-DETECTION OVERLAPPING INFERENCE FOR END-TO-END LONG-FORM SPEECH RECOGNITION	6977
<i>Jinhan Wang, University of California, Los Angeles, United States of America; Xiaosu Tong, Jinxi Guo, Di He, Roland Maas, Amazon, United States of America</i>	
SPE-34.3: TORCHAUDIO: BUILDING BLOCKS FOR AUDIO AND SPEECH PROCESSING	6982
<i>Yao-Yuan Yang, University of California, San Diego, United States of America; Moto Hira, Zhaoheng Ni, Artyom Astafurov, Caroline Chen, Christian Puhersch, Dmitriy Genzel, Donny Greenberg, Edward Yang, Jason Lian, Jeff Hwang, Ji Chen, Soumith Chintala, Vincent Quenneville-Bélair, Facebook AI, United States of America; David Pollack, Solvemate GmbH, Germany; Peter Goldsborough, Anduril Industries, United States of America; Sean Narenthiran, Grid AI Labs, United Kingdom of Great Britain and Northern Ireland; Shinji Watanabe, Carnegie Mellon University, United States of America</i>	
SPE-34.4: UNSUPERVISED MODEL ADAPTATION FOR END-TO-END ASR	6987
<i>Ganesh Sivaraman, Ricardo Casal, Matthew Garland, Elie Khoury, Pindrop, United States of America</i>	
SPE-34.5: SPEECH RECOGNITION USING BIOLOGICALLY-INSPIRED NEURAL NETWORKS	6992
<i>Thomas Bohnstingl, Ayush Garg, Stanisław Woźniak, Evangelos Eleftheriou, Angeliki Pantazi, IBM Research Zurich, Switzerland; George Saon, IBM Research AI, United States of America</i>	

SPE-35: VOICE CONVERSION II

SPE-35.1: ASSEM-VC: REALISTIC VOICE CONVERSION BY ASSEMBLING MODERN 6997 SPEECH SYNTHESIS TECHNIQUES

Kang-wook Kim, Junhyeok Lee, Myun-chul Joe, MINDsLab Inc., Korea, Republic of; Seung-won Park, Seoul National University, Korea, Republic of

SPE-35.2: MINIMIZING RESIDUALS FOR NATIVE-NONNATIVE VOICE 7002 CONVERSION IN A SPARSE, ANCHOR-BASED REPRESENTATION OF SPEECH

Christopher Liberatore, Ricardo Gutierrez-Osuna, Texas A&M University, United States of America

SPE-35.3: IMPROVING RECOGNITION-SYNTHESIS BASED ANY-TO-ONE VOICE 7007 CONVERSION WITH CYCLIC TRAINING

Yannian Chen, Zhenhua Ling, National Engineering Laboratory for Speech and Language Information Processing, University of Science and Technology of China, Hefei, P.R.China, China; Lijuan Liu, Yajun Hu, Yuan Jiang, iFLYTEK Research, iFLYTEK Co., Ltd., Hefei, P.R.China, China

SPE-35.4: NVC-NET: END-TO-END ADVERSARIAL VOICE CONVERSION..... 7012

Bac Nguyen, Fabien Cardinaux, Sony Europe B.V., Germany

SPE-35.5: U-GAT-VC: UNSUPERVISED GENERATIVE ATTENTIONAL NETWORKS 7017 FOR NON-PARALLEL VOICE CONVERSION

Sheng Shi, Yangzhou Du, Jianping Fan, Lenovo Research, China; Jiahao Shao, Tsinghua University, China; Yifei Hao, University of Southern California, China

SPE-35.6: DISENTANGLING CONTENT AND FINE-GRAINED PROSODY 7022 INFORMATION VIA HYBRID ASR BOTTLENECK FEATURES FOR VOICE CONVERSION

Xintao Zhao, Changhe Song, Zhiyong Wu, Tsinghua University, China; Feng Liu, Shiyin Kang, Deyi Tuo, Huya Inc, China; Helen Meng, The Chinese University of Hong Kong, China

SPE-36: DIALOG SYSTEM III: EMOTION RECOGNITION AND PARALINGUISTIC MODELING

SPE-36.1: A COMMONSENSE KNOWLEDGE ENHANCED NETWORK WITH 7027 RETROSPECTIVE LOSS FOR EMOTION RECOGNITION IN SPOKEN DIALOG

Yunhe Xie, Chengjie Sun, Zhenzhou Ji, Harbin Institute of Technology, China

SPE-36.2: HIERARCHICAL AND MULTI-VIEW DEPENDENCY MODELLING 7032 NETWORK FOR CONVERSATIONAL EMOTION RECOGNITION

Yu-Ping Ruan, Shu-Kai Zheng, Taihao Li, Fen Wang, Guanxiong Pei, Zhejiang Lab, China

SPE-36.3: MM-DFN: MULTIMODAL DYNAMIC FUSION NETWORK FOR EMOTION 7037 RECOGNITION IN CONVERSATIONS

Dou Hu, Xiaolong Hou, Lianxin Jiang, Ping An Life Insurance Company of China, Ltd., China; Lingwei Wei, Institute of Information Engineering, Chinese Academy of Sciences, China; Yang Mo, Ping An Life Insurance Company of China, Ltd, China

SPE-36.4: MODELING INTENTION, EMOTION AND EXTERNAL WORLD IN 7042 DIALOGUE SYSTEMS

Wei Peng, Yue Hu, Luxi Xing, Yuqiang Xie, Xingsheng Zhang, Yajing Sun, Institute of Information Engineering, Chinese Academy of Sciences, China

SPE-36.5: A NEURAL PROSODY ENCODER FOR END-TO-END DIALOGUE ACT 7047 CLASSIFICATION

Kai Wei, Martin Radfar, Thanh Tran, Markus Mueller, Grant Strimel, Nathan Susanj, Athanasios Mouchtaris, Maurizio Omologo, Alexa Speech, Amazon, United States of America; Dillon Knox, University of Southern California, United States of America

SPE-36.6: IMPROVING CONTEXTUAL COHERENCE IN VARIATIONAL PERSONALIZED AND EMPATHETIC DIALOGUE AGENTS	7052
<i>Jing Yang Lee, Woon Seng Gan, Nanyang Technological University, Singapore; Kong Aik Lee, A*STAR, Singapore</i>	
SPE-37: SPEAKER RECOGNITION V: APPLICATIONS	
SPE-37.1: FUSION AND ORTHOGONAL PROJECTION FOR IMPROVED FACE-VOICE ASSOCIATION	7057
<i>Muhammad Saad Saeed, Muhammad Haroon Yousaf, University of Engineering and Technology Taxila, Pakistan; Muhammad Haris Khan, Mohamed Bin Zayed University of Artificial Intelligence, United Arab Emirates; Shah Nawaz, Alessio Del Bue, Istituto Italiano di Tecnologia, Italy</i>	
SPE-37.2: OPENFEAT: IMPROVING SPEAKER IDENTIFICATION BY OPEN-SET FEW-SHOT EMBEDDING ADAPTATION WITH TRANSFORMER	7062
<i>Kishan K C, Rochester Institute of Technology, United States of America; Zhenning Tan, Long Chen, Minh Jin, Eunjung Han, Andreas Stolcke, Chul Lee, Amazon Alexa AI, United States of America</i>	
SPE-37.3: TOWARDS LIGHTWEIGHT APPLICATIONS: ASYMMETRIC ENROLL-VERIFY STRUCTURE FOR SPEAKER VERIFICATION	7067
<i>Qingjian Lin, Lin Yang, Xuyang Wang, Junjie Wang, Lenovo, China; Xiaoyi Qin, Ming Li, Duke Kunshan University, China</i>	
SPE-37.4: SPEAKER EMBEDDING CONVERSION FOR BACKWARD AND CROSS-CHANNEL COMPATIBILITY	7072
<i>Tianxiang Chen, Elie Khoury, Pindrop, United States of America</i>	
SPE-37.5: IMPROVING FAIRNESS IN SPEAKER VERIFICATION VIA GROUP-ADAPTED FUSION NETWORK	7077
<i>Hua Shen, The Pennsylvania State University, United States of America; Yuguang Yang, Guoli Sun, Ryan Langman, Eunjung Han, Jasha Droppo, Andreas Stolcke, Amazon, United States of America</i>	
SPE-37.6: CS-REP: MAKING SPEAKER VERIFICATION NETWORKS EMBRACING RE-PARAMETERIZATION	7082
<i>Ruiteng Zhang, Wenhuan Lu, Junhai Xu, Tianjin University, China; Jianguo Wei, Tianjin University; Qinghai Nationalities University, China; Lin Zhang, National Institute of Informatics, Japan; Yantao Ji, Xi'an Jiaotong University, China; Xugang Lu, National Institute of Information and Communications Technology, Japan</i>	
SPE-38: SELF SUPERVISED LEARNING FOR SPEECH RECOGNITION	
SPE-38.1: DISTILHUBERT: SPEECH REPRESENTATION LEARNING BY LAYER-WISE DISTILLATION OF HIDDEN-UNIT BERT	7087
<i>Heng-Jui Chang, Shu-wen Yang, Hung-yi Lee, National Taiwan University, Taiwan</i>	
SPE-38.2: IMPROVING SELF-SUPERVISED LEARNING FOR SPEECH RECOGNITION WITH INTERMEDIATE LAYER SUPERVISION	7092
<i>Chengyi Wang, Zhenglu Yang, Nankai University, China; Yu Wu, Sanyuan Chen, Shujie Liu, Jinyu Li, Yao Qian, Microsoft, China</i>	
SPE-38.3: WAV2VEC-SWITCH: CONTRASTIVE LEARNING FROM ORIGINAL-NOISY SPEECH PAIRS FOR ROBUST SPEECH RECOGNITION	7097
<i>Yiming Wang, Jinyu Li, Yao Qian, Chengyi Wang, Yu Wu, Microsoft Corporation, United States of America; Heming Wang, The Ohio State University, United States of America</i>	
SPE-38.4: EFFICIENT ADAPTER TRANSFER OF SELF-SUPERVISED SPEECH MODELS FOR AUTOMATIC SPEECH RECOGNITION	7102
<i>Bethan Thomas, Salah Karout, Huawei R&D UK, United Kingdom of Great Britain and Northern Ireland; Samuel Kessler, University of Oxford, United Kingdom of Great Britain and Northern Ireland</i>	

SPE-38.5: AN EXPLORATION OF HUBERT WITH LARGE NUMBER OF CLUSTER 7107
UNITS AND MODEL ASSESSMENT USING BAYESIAN INFORMATION CRITERION
Takashi Maekaku, Yuya Fujita, Yahoo Japan Corporation, Japan; Xuankai Chang, Shinji Watanabe, Carnegie Mellon University, United States of America

SPE-38.6: OPTIMIZE WAV2VEC2S ARCHITECTURE FOR SMALL TRAINING SET 7112
THROUGH ANALYZING ITS PRE-TRAINED MODELS ATTENTION PATTERN
Liu Chen, Meysam Asgari, Hiroko Dodge, Oregon Health & Science University, United States of America

SPE-39: SPEECH SYNTHESIS: FRONT-END

SPE-39.1: PART-OF-SPEECH MODELS COMPRESSION METHODS FOR ON-DEVICE 7117
GRAPHEME-TO-PHONEME CONVERSION
Marek Kubis, Paweł Skórzewski, Adam Mickiewicz University in Poznań, Poland; Maxime Méloux, Marcin Lewandowski, Samsung R&D Institute Poland, Poland; Gunu Jho, Hyoungmin Park, Samsung Electronics, Korea, Republic of

SPE-39.2: AN END-TO-END CHINESE TEXT NORMALIZATION MODEL BASED ON 7122
RULE-GUIDED FLAT-LATTICE TRANSFORMER
Wenlin Dai, Changhe Song, Xiang Li, Zhiyong Wu, Tsinghua University, China; Huashan Pan, Xiulin Li, Databaker Technology Co., Ltd, China; Helen Meng, The Chinese University of Hong Kong, China

SPE-39.3: CHINESE SPELLING TEXT GENERATION OF MATHEMATICAL FORMULAS 7127
Su Dong, Sicen Liu, Buzhou Tang, Harbin Institute of Technology, China; Shan Liu, Tencent Technology Company, China

SPE-39.4: POLYPHONE DISAMBIGUATION AND ACCENT PREDICTION USING 7132
PRE-TRAINED LANGUAGE MODELS IN JAPANESE TTS FRONT-END
Rem Hida, Masaki Hamada, Emiru Tsunoo, Toshiyuki Sekiya, Sony Group Corporation, Japan; Chie Kamada, Toshiyuki Kumakura, Sony Corporation of America, United States of America

SPE-39.6: DATA AUGMENTATION FOR LONG-TAILED AND IMBALANCED 7137
POLYPHONE DISAMBIGUATION IN MANDARIN
Yang Zhang, Haitong Zhang, Yue Lin, NetEase Games AI Lab, China

SPE-40: LANGUAGE UNDERSTANDING I: END-TO-END FRAMEWORK

SPE-40.1: LEVERAGING BILINEAR ATTENTION TO IMPROVE SPOKEN LANGUAGE 7142
UNDERSTANDING
Dongsheng Chen, Zhiqi Huang, Yuexian Zou, Peking University, China

SPE-40.2: BUILDING ROBUST SPOKEN LANGUAGE UNDERSTANDING BY CROSS 7147
ATTENTION BETWEEN PHONEME SEQUENCE AND ASR HYPOTHESIS
Zexun Wang, Yuming Zhao, Mingchao Feng, Meng Chen, Xiaodong He, JD AI, China; Yuquan Le, Hunan University, China; Yi Zhu, University of Cambridge, United Kingdom of Great Britain and Northern Ireland

SPE-40.3: INTEGRATION OF PRE-TRAINED NETWORKS WITH CONTINUOUS 7152
TOKEN INTERFACE FOR END-TO-END SPOKEN LANGUAGE UNDERSTANDING
Seunghyun Seo, Donghyun Kwak, Naver Corporation, Korea, Republic of; Bowon Lee, Inha University, Korea, Republic of

SPE-40.4: TIE YOUR EMBEDDINGS DOWN: CROSS-MODAL LATENT SPACES FOR 7157
END-TO-END SPOKEN LANGUAGE UNDERSTANDING
Bhuvan Agrawal, Markus Müller, Samridhi Choudhary, Martin Radfar, Athanasios Mouchtaris, Ross McGowan, Nathan Susanj, Siegfried Kunzmann, Amazon, United States of America

SPE-40.5: IMPROVING END-TO-END MODELS FOR SET PREDICTION IN SPOKEN LANGUAGE UNDERSTANDING	7162
<i>Hong-Kwang Kuo, Zoltan Tuske, Samuel Thomas, Brian Kingsbury, George Saon, IBM Research AI, United States of America</i>	
SPE-40.6: ESPNET-SLU: ADVANCING SPOKEN LANGUAGE UNDERSTANDING THROUGH ESPNET	7167
<i>Siddhant Arora, Siddharth Dalmia, Xuankai Chang, Yushi Ueda, Yifan Peng, Sujay Kumar, Karthik Ganesan, Brian Yan, Alan W Black, Shinji Watanabe, Carnegie Mellon University, United States of America; Pavel Denisov, Ngoc Thang Vu, University of Stuttgart, Germany; Yuekai Zhang, Zoom Video Communications, China</i>	
SPE-41: SPEAKER RECOGNITION VI: DOMAIN ADAPTATION	
SPE-41.1: THE CORAL++ ALGORITHM FOR UNSUPERVISED DOMAIN ADAPTATION OF SPEAKER RECOGNITION	7172
<i>Rongjin Li, Weibin Zhang, Dongpeng Chen, VoiceAI Technologies, Co. Ltd. Shenzhen, China, China</i>	
SPE-41.2: LEARNING DOMAIN-INVARIANT TRANSFORMATION FOR SPEAKER VERIFICATION	7177
<i>Hanyi Zhang, Longbiao Wang, Meng Liu, Jianwu Dang, Hui Chen, Tianjin University, China; Kong Aik Lee, A*STAR, Singapore</i>	
SPE-41.3: DOMAIN ROBUST DEEP EMBEDDING LEARNING FOR SPEAKER RECOGNITION	7182
<i>Hang-Rui Hu, Yan Song, Ying Liu, Li-Rong Dai, University of Science and Technology of China, China; Ian McLoughlin, Singapore Institute of Technology, Singapore; Lin Liu, iFLYTEK CO. LTD., China</i>	
SPE-41.4: TACKLING THE SCORE SHIFT IN CROSS-LINGUAL SPEAKER VERIFICATION BY EXPLOITING LANGUAGE INFORMATION	7187
<i>Jenthe Thienpondt, Brecht Desplanques, Kris Demuyne, Ghent University, Belgium</i>	
SPE-41.5: DOMAIN ADAPTATION FOR SPEAKER RECOGNITION IN SINGING AND SPOKEN VOICE	7192
<i>Anurag Chowdhury, Austin Cozzo, Arun Ross, Michigan State University, United States of America</i>	
SPE-41.6: CDMA: CROSS-DOMAIN DISTANCE METRIC ADAPTATION FOR SPEAKER VERIFICATION	7197
<i>Jianchen Li, Jiqing Han, Hongwei Song, Harbin Institute of Technology, China</i>	
SPE-42: SPEECH RECOGNITION: ACOUSTIC MODELING I	
SPE-42.1: WORD ORDER DOES NOT MATTER FOR SPEECH RECOGNITION	7202
<i>Vineel Pratap Konduru, Qiantong Xu, Tatiana Likhomanenko, Gabriel Synnaeve, Ronan Collobert, Facebook AI, United States of America</i>	
SPE-42.2: CONTRASTIVE SIAMESE NETWORK FOR SEMI-SUPERVISED SPEECH RECOGNITION	7207
<i>Soheil Khorram, Jaeyoung Kim, Anshuman Tripathi, Han Lu, Qian Zhang, Hasim Sak, Google Inc. USA, United States of America</i>	
SPE-42.3: SEQUENCE TRANSDUCTION WITH GRAPH-BASED SUPERVISION	7212
<i>Niko Moritz, Facebook AI, United Kingdom of Great Britain and Northern Ireland; Takaaki Hori, Jonathan Le Roux, MERL, United States of America; Shinji Watanabe, Carnegie Mellon University, United States of America</i>	
SPE-42.4: SPEECHMOE2: MIXTURE-OF-EXPERTS MODEL WITH IMPROVED ROUTING	7217
<i>Zhao You, Shulin Feng, Dan Su, Dong Yu, Tencent, China</i>	

SPE-42.5: SUPERVISED ATTENTION IN SEQUENCE-TO-SEQUENCE MODELS FOR 7222
SPEECH RECOGNITION

Gene-Ping Yang, Hao Tang, University of Edinburgh, United Kingdom of Great Britain and Northern Ireland

SPE-42.6: END-TO-END SPEECH RECOGNITION FROM FEDERATED ACOUSTIC 7227
MODELS

Yan Gao, Pedro P. B. de Gusmao, Nicholas D. Lane, University of Cambridge, United Kingdom of Great Britain and Northern Ireland; Titouan Parcollet, Avignon University, France; Salah Zaiem, Telecom Paris, France; Javier Fernandez-Marques, University of Oxford, United Kingdom of Great Britain and Northern Ireland; Daniel J. Beutel, Adap GmbH; University of Cambridge, Germany

SPE-43: SPEECH SYNTHESIS: SINGING VOICE AND OTHERS

SPE-43.1: HIFIDENOISE: HIGH-FIDELITY DENOISING TEXT TO SPEECH WITH 7232
ADVERSARIAL NETWORKS

Lichao Zhang, Yi Ren, Zhou Zhao, Zhejiang University, China; Liqun Deng, Huawei Noah's Ark Lab, China

SPE-43.2: VISINGER: VARIATIONAL INFERENCE WITH ADVERSARIAL LEARNING 7237
FOR END-TO-END SINGING VOICE SYNTHESIS

Yongmao Zhang, Jian Cong, Heyang Xue, Lei Xie, Northwestern Polytechnical University, China; Pengcheng Zhu, Mengxiao Bi, Fuxi AI Lab, NetEase Inc., China

SPE-43.3: A MELODY-UNSUPERVISION MODEL FOR SINGING VOICE SYNTHESIS 7242

Soonbeom Choi, Juhan Nam, Korea Advanced Institute of Science and Technology, Korea, Republic of

SPE-43.4: TRANSFORMER-S2A: ROBUST AND EFFICIENT SPEECH-TO-ANIMATION 7247

Liyang Chen, Zhiyong Wu, Shenzhen International Graduate School, China; Jun Ling, Shanghai Jiao Tong University, China; Runnan Li, Xu Tan, Sheng Zhao, Microsoft, China, China

SPE-43.5: VCVTS: MULTI-SPEAKER VIDEO-TO-SPEECH SYNTHESIS VIA 7252
CROSS-MODAL KNOWLEDGE TRANSFER FROM VOICE CONVERSION

Disong Wang, Xunying Liu, Helen Meng, The Chinese University of Hong Kong, Hong Kong; Shan Yang, Dan Su, Dong Yu, Tencent AI Lab, China

SPE-44: LANGUAGE ASSESSMENT / TRANSCRIPTION AUGMENTATION

SPE-44.1: FAST TASK-SPECIFIC ADAPTATION IN SPOKEN LANGUAGE ASSESSMENT 7257
WITH META-LEARNING

Binghuai Lin, Liyuan Wang, Tencent Technology Co., Ltd, China

SPE-44.2: TRANSFORMER-BASED MULTI-ASPECT MULTI-GRANULARITY 7262
NON-NATIVE ENGLISH SPEAKER PRONUNCIATION ASSESSMENT

Yuan Gong, James Glass, Massachusetts Institute of Technology, United States of America; Ziyi Chen, Iek-Heng Chu, Peng Chang, PAII Inc., United States of America

SPE-44.3: A MODEL FOR ASSESSOR BIAS IN AUTOMATIC PRONUNCIATION 7267
ASSESSMENT

Jose Antonio Lopez Saenz, Thomas Hain, University of Sheffield, United Kingdom of Great Britain and Northern Ireland

SPE-44.4: UNIFIED MULTIMODAL PUNCTUATION RESTORATION FRAMEWORK 7272
FOR MIXED-MODALITY CORPUS

Yaoming Zhu, Liwei Wu, Shanbo Cheng, Mingxuan Wang, Bytedance, China

SPE-44.5: PUNCTUATION PREDICTION FOR STREAMING ON-DEVICE SPEECH 7277
RECOGNITION

Zhikai Zhou, Yanmin Qian, Shanghai Jiao Tong University, China; Tian Tan, AISpeech Ltd., China

SPE-44.6: ASR ERROR CORRECTION WITH DUAL-CHANNEL SELF-SUPERVISED LEARNING	7282
<i>Fan Zhang, Jinyao Yan, Communication University of China, China; Mei Tu, Song Liu, Samsung Research China - Beijing (SRCB), China</i>	
SPE-45: SPEECH ENHANCEMENT: SPEECH EXTRACTION AND AUDIO-VISUAL ENHANCEMENT	
SPE-45.1: L-SPEX: LOCALIZED TARGET SPEAKER EXTRACTION	7287
<i>Meng Ge, Longbiao Wang, Jianwu Dang, Tianjin University, China; Chenglin Xu, Kuaishou Technology, China; Eng Siong Chng, Nanyang Technological University, Singapore; Haizhou Li, National University of Singapore, Singapore</i>	
SPE-45.3: DPCCN: DENSELY-CONNECTED PYRAMID COMPLEX CONVOLUTIONAL NETWORK FOR ROBUST SPEECH SEPARATION AND EXTRACTION	7292
<i>Jiangyu Han, Yanhua Long, Shanghai Normal University, China; Lukas Burget, Jan Cernocky, Brno University of Technology, Czechia</i>	
SPE-45.4: MIXED PRECISION DNN QUANTIZATION FOR OVERLAPPED SPEECH SEPARATION AND RECOGNITION	7297
<i>Junhao Xu, Xunying Liu, Helen Mei-Ling Meng, The Chinese University of Hong Kong, China; Jianwei Yu, Tencent AI Lab, China</i>	
SPE-45.5: THE IMPACT OF REMOVING HEAD MOVEMENTS ON AUDIO-VISUAL SPEECH ENHANCEMENT	7302
<i>Zhiqi Kang, Radu Horaud, Xavier Alameda-Pineda, Inria Grenoble Rhône-Alpes & Univ. Grenoble Alpes, France, France; Mostafa Sadeghi, Inria Nancy Grand-Est, France, France; Jacob Donley, Anurag Kumar, Facebook Reality Labs Research, Redmond WA, USA, United States of America</i>	
SPE-45.6: VSEGAN: VISUAL SPEECH ENHANCEMENT GENERATIVE ADVERSARIAL NETWORK	7307
<i>Xinmeng Xu, Trinity College Dublin, China; Yang Wang, Dongxiang Xu, Yiyuan Peng, Cong Zhang, Jie Jia, Binbin Chen, Vivo Communication Technology Co. Ltd., China</i>	
SPE-46: MULTI SPEAKER AND MULTI-CHANNEL SPEECH RECOGNITION AND PROCESSING	
SPE-46.1: ENDPOINT DETECTION FOR STREAMING END-TO-END MULTI-TALKER ASR	7312
<i>Liang Lu, Otter.ai, United States of America; Jinyu Li, Yifan Gong, Microsoft, USA, United States of America</i>	
SPE-46.2: CONTINUOUS STREAMING MULTI-TALKER ASR WITH DUAL-PATH TRANSDUCERS	7317
<i>Desh Raj, Johns Hopkins University, United States of America; Liang Lu, Zhuo Chen, Yashesh Gaur, Jinyu Li, Microsoft Corp., United States of America</i>	
SPE-46.3: EXTENDED GRAPH TEMPORAL CLASSIFICATION FOR MULTI-SPEAKER END-TO-END ASR	7322
<i>Xuankai Chang, Shinji Watanabe, Carnegie Mellon University, United States of America; Niko Moritz, Facebook AI, United Kingdom of Great Britain and Northern Ireland; Takaaki Hori, Jonathan Le Roux, Mitsubishi Electric Research Laboratories (MERL), United States of America</i>	
SPE-46.4: ADA-VAD: UNPAIRED ADVERSARIAL DOMAIN ADAPTATION FOR NOISE-ROBUST VOICE ACTIVITY DETECTION	7327
<i>Taesoo Kim, Jong Hwan Ko, Sungkyunkwan university, Korea, Republic of; Jiho Chang, Korea Research Institute of Standards and Science, Korea, Republic of</i>	

SPE-46.5: MULTI-CHANNEL END-TO-END NEURAL DIARIZATION WITH DISTRIBUTED MICROPHONES	7332
<i>Shota Horiguchi, Yuki Takashima, Yohei Kawaguchi, Hitachi, Ltd., Japan; Paola Garcia, Johns Hopkins University, United States of America; Shinji Watanabe, Carnegie Mellon University, United States of America</i>	
SPE-46.6: MULTI-CHANNEL SPEAKER DIARIZATION USING SPATIAL FEATURES FOR MEETINGS	7337
<i>Naijun Zheng, XunYing Liu, Helen Meng, The Chinese University of Hong Kong, Hong Kong; Na Li, Jianwei Yu, Chao Weng, Dan Su, Tencent AI Lab, Hong Kong</i>	
SPE-47: EMOTION RECOGNITION: GENERAL TOPICS I	
SPE-47.1: SPEAKER NORMALIZATION FOR SELF-SUPERVISED SPEECH EMOTION RECOGNITION	7342
<i>Itai Gat, Hagai Aronowitz, Weizhong Zhu, Edmilson Morais, Ron Hoory, IBM Research, Israel</i>	
SPE-47.2: SENTIMENT-AWARE AUTOMATIC SPEECH RECOGNITION PRE-TRAINING FOR ENHANCED SPEECH EMOTION RECOGNITION	7347
<i>Ayoub Ghriss, University of Colorado, Boulder, United States of America; Bo Yang, Viktor Rozgic, Wang Chao, Amazon LLC, United States of America; Elizabeth Shriberg, Ellipsis Health, United States of America</i>	
SPE-47.3: CONFIDENCE ESTIMATION FOR SPEECH EMOTION RECOGNITION BASED ON THE RELATIONSHIP BETWEEN EMOTION CATEGORIES AND PRIMITIVES	7352
<i>Yang Li, University of Rochester, United States of America; Constantinos Papayiannis, Viktor Rozgic, Elizabeth Shriberg, Chao Wang, Amazon, United States of America</i>	
SPE-47.4: AUXFORMER: ROBUST APPROACH TO AUDIOVISUAL EMOTION RECOGNITION	7357
<i>Lucas Goncalves, Carlos Busso, The University of Texas at Dallas, United States of America</i>	
SPE-47.5: FUSING ASR OUTPUTS IN JOINT TRAINING FOR SPEECH EMOTION RECOGNITION	7362
<i>Yuanhao Li, Peter Bell, Catherine Lai, University of Edinburgh, United Kingdom of Great Britain and Northern Ireland</i>	
SPE-47.6: SPEECH EMOTION RECOGNITION WITH CO-ATTENTION BASED MULTI-LEVEL ACOUSTIC INFORMATION	7367
<i>Heqing Zou, Chen Chen, Deepu Rajan, Eng Siong Chng, Nanyang Technological University, Singapore; Yuke Si, Tianjin University, China</i>	
SPE-48: LANGUAGE DISORDERS: RECOGNITION	
SPE-48.1: MULTI-MODAL ACOUSTIC-ARTICULATORY FEATURE FUSION FOR DYSARTHIC SPEECH RECOGNITION	7372
<i>Zhengjun Yue, Heidi Christensen, Jon Barker, The University of Sheffield, United Kingdom of Great Britain and Northern Ireland; Erfan Loweimi, Zoran Cvetkovic, King's College London, United Kingdom of Great Britain and Northern Ireland</i>	
SPE-48.2: RAW SOURCE AND FILTER MODELLING FOR DYSARTHIC SPEECH RECOGNITION	7377
<i>Zhengjun Yue, The University of Sheffield, United Kingdom of Great Britain and Northern Ireland; Erfan Loweimi, Zoran Cvetkovic, King's College London, United Kingdom of Great Britain and Northern Ireland</i>	
SPE-48.3: SYNTHESIZING DYSARTHIC SPEECH USING MULTI-SPEAKER TTS FOR DYSARTHIC SPEECH RECOGNITION	7382
<i>Mohammad Soleymanpour, Michael T. Johnson, University of Kentucky, United States of America; Rahim Soleymanpour, University of Connecticut, United States of America; Jeffrey Berry, Marquette University, United States of America</i>	

SPE-48.4: TOWARDS INTERPRETING DEEP LEARNING MODELS TO UNDERSTAND LOSS OF SPEECH INTELLIGIBILITY IN SPEECH DISORDERS STEP 2: CONTRIBUTION OF THE EMERGENCE OF PHONETIC TRAITS	7387
<i>Sondes Abderrazek, Corinne Fredouille, LIA - Avignon Université, France; Alain Ghio, Muriel Lalain, Christine Meunier, LPL, Aix-Marseille Université, France; Virginie Woisard, UT2J, Octogone-Lordat, Toulouse University, Toulouse Hospital, France</i>	
SPE-48.5: CONSTANT Q CEPSTRAL COEFFICIENTS FOR CLASSIFICATION OF NORMAL VS. PATHOLOGICAL INFANT CRY	7392
<i>Hemant A. Patil, Ankur T. Patil, Aastha Kachhi, Dhirubhai Ambani Institute of Information and Communication Technology, India</i>	
SPE-48.6: NONVERBAL SOUND DETECTION FOR DISORDERED SPEECH	7397
<i>Colin Lea, Zifang Huang, Lauren Tooley, Zeinab Liaghat, Shri Thelapurath, Leah Findlater, Jeffrey P. Bigham, Apple, United States of America; Dhruv Jain, University of Washington, United States of America</i>	
SPE-49: SPEECH ENHANCEMENT AND DEREVERBERATION	
SPE-49.1: CONDITIONAL DIFFUSION PROBABILISTIC MODEL FOR SPEECH ENHANCEMENT	7402
<i>Yen-Ju Lu, Zhong-Qiu Wang, Shinji Watanabe, Carnegie Mellon University, United States of America; Alexander Richard, Facebook, United States of America; Cheng Yu, Yu Tsao, Academia Sinica, Taiwan</i>	
SPE-49.2: DEEPFILTERNET: A LOW COMPLEXITY SPEECH ENHANCEMENT FRAMEWORK FOR FULL-BAND AUDIO BASED ON DEEP FILTERING	7407
<i>Hendrik Schröter, Andreas Maier, Friedrich-Alexander University Erlangen-Nuremberg (FAU), Germany; Alberto N. Escalante-B., Tobias Rosenkranz, WS Audiology, Germany</i>	
SPE-49.3: METRICGAN-U: UNSUPERVISED SPEECH ENHANCEMENT/ DEREVERBERATION BASED ONLY ON NOISY/ REVERBERATED SPEECH	7412
<i>Szu-Wei Fu, Microsoft, United States of America; Cheng Yu, Kuo-Hsuan Hung, Yu Tsao, Academia Sinica, United States of America; Mirco Ravanelli, Mila-Quebec AI Institute, United States of America</i>	
SPE-49.4: UFORMER: A UNET BASED DILATED COMPLEX & REAL DUAL-PATH CONFORMER NETWORK FOR SIMULTANEOUS SPEECH ENHANCEMENT AND DEREVERBERATION	7417
<i>Yihui Fu, Shubo Lv, Yukai Jv, Lei Xie, Audio, Speech and Language Processing Group (ASLP@NPU), Northwestern Polytechnical University, Xi'an, China, China; Yun Liu, Jingdong Li, Dawei Luo, AI Interaction Division, Sogou Inc., Beijing, China, China</i>	
SPE-49.5: ATTENUATION OF ACOUSTIC EARLY REFLECTIONS IN TELEVISION STUDIOS USING PRETRAINED SPEECH SYNTHESIS NEURAL NETWORK	7422
<i>Tomer Rosenbaum, Israel Cohen, Technion – Israel Institute of Technology, Israel; Emil Winebrand, Insoundz Ltd., Israel</i>	
SPE-50: SPEECH RECOGNITION: ACOUSTIC MODELING II	
SPE-50.1: NON-AUTOREGRESSIVE ASR WITH SELF-CONDITIONED FOLDED ENCODERS	7427
<i>Tatsuya Komatsu, LINE Corporation, Japan</i>	
SPE-50.2: HAVE BEST OF BOTH WORLDS: TWO-PASS HYBRID AND E2E CASCADING FRAMEWORK FOR SPEECH RECOGNITION	7432
<i>Guoli Ye, Vadim Mazalov, Jinyu Li, Yifan Gong, Microsoft, United States of America</i>	

SPE-50.3: CONFORMER-BASED HYBRID ASR SYSTEM FOR SWITCHBOARD DATASET	7437
<i>Mohammad Zeineldeen, Christoph Lüscher, Wilfried Michel, Ralf Schlüter, Hermann Ney, RWTH Aachen University / AppTek, Germany; Jingjing Xu, Alexander Gerstenberger, RWTH Aachen University, Germany</i>	
SPE-50.4: IMPROVING FACTORED HYBRID HMM ACOUSTIC MODELING WITHOUT STATE TYING	7442
<i>Tina Raissi, Ralf Schlüter, RWTH Aachen University, Germany; Eugen Beck, Apptek GmbH, Germany; Hermann Ney, RWTH Aachen University, Germany</i>	
SPE-50.5: AUDITORY-BASED DATA AUGMENTATION FOR END-TO-END AUTOMATIC SPEECH RECOGNITION	7447
<i>Zehai Tu, Jack Deadman, Ning Ma, Jon Barker, University of Sheffield, United Kingdom of Great Britain and Northern Ireland</i>	
SPE-50.6: DELIBERATION OF STREAMING RNN-TRANSDUCER BY NON-AUTOREGRESSIVE DECODING	7452
<i>Weiran Wang, Ke Hu, Tara Sainath, Google, United States of America</i>	
SPE-51: SPEECH SYNTHESIS: NOVEL ACOUSTIC MODELS	
SPE-51.1: NEURAL HMMS ARE ALL YOU NEED (FOR HIGH-QUALITY ATTENTION-FREE TTS)	7457
<i>Shivam Mehta, Éva Székely, Jonas Beskow, Gustav Eje Henter, KTH Royal Institute of Technology, Sweden</i>	
SPE-51.2: AUTOREGRESSIVE VARIATIONAL AUTOENCODER WITH A HIDDEN SEMI-MARKOV MODEL-BASED STRUCTURED ATTENTION FOR SPEECH SYNTHESIS	7462
<i>Takato Fujimoto, Kei Hashimoto, Yoshihiko Nankaku, Keichi Tokuda, Nagoya Institute of Technology, Japan</i>	
SPE-51.3: PAMA-TTS: PROGRESSION-AWARE MONOTONIC ATTENTION FOR STABLE SEQ2SEQ TTS WITH ACCURATE PHONEME DURATION CONTROL	7467
<i>Yunchao He, Jian Luan, Yujun Wang, Xiaomi Corporation, China</i>	
SPE-51.4: IMPROVING FASTSPEECH TTS WITH EFFICIENT SELF-ATTENTION AND COMPACT FEED-FORWARD NETWORK	7472
<i>Yujia Xiao, Xi Wang, Lei He, Frank K. Soong, Microsoft, China</i>	
SPE-51.5: VARIANCEFLOW: HIGH-QUALITY AND CONTROLLABLE TEXT-TO-SPEECH USING VARIANCE INFORMATION VIA NORMALIZING FLOW	7477
<i>Yoonhyung Lee, Kyomin Jung, Seoul National University, Korea, Republic of; Jinhyeok Yang, NCSoft, Korea, Republic of</i>	
SPE-51.6: MIXER-TTS: NON-AUTOREGRESSIVE, FAST AND COMPACT TEXT-TO-SPEECH MODEL CONDITIONED ON LANGUAGE MODEL EMBEDDINGS	7482
<i>Oktai Tatanov, Stanislav Beliaev, Boris Ginsburg, NVIDIA, Russian Federation</i>	
SPE-52: LANGUAGE UNDERSTANDING II: CONTEXT/KNOWLEDGE INTEGRATION	
SPE-52.1: KNOWLEDGE AUGMENTED BERT MUTUAL NETWORK IN MULTI-TURN SPOKEN DIALOGUES	7487
<i>Ting-Wei Wu, Biing-Hwang Juang, Georgia Institute of Technology, United States of America</i>	
SPE-52.2: TINYS2I: A SMALL-FOOTPRINT UTTERANCE CLASSIFICATION MODEL WITH CONTEXTUAL SUPPORT FOR ON-DEVICE SLU	7492
<i>Anastasios Alexandridis, Kanthashree Mysore Sathyendra, Grant Strimel, Pavel Kveton, Jon Webb, Athanasios Mouchtaris, Amazon.com, United States of America</i>	

SPE-52.3: TOWARDS END-TO-END INTEGRATION OF DIALOG HISTORY FOR IMPROVED SPOKEN LANGUAGE UNDERSTANDING	7497
<i>Vishal Sunder, Eric Fosler-Lussier, The Ohio State University, United States of America; Samuel Thomas, Hong-Kwang Kuo, Jatin Ganhotra, Brian Kingsbury, IBM, United States of America</i>	
SPE-52.4: IMPROVING SPOKEN LANGUAGE UNDERSTANDING BY ENHANCING TEXT REPRESENTATION	7502
<i>Thai Binh Nguyen, Karlsruhe Institute of Technology, Germany; , ,</i>	
SPE-52.5: MULTI-TASK RNN-T WITH SEMANTIC DECODER FOR STREAMABLE SPOKEN LANGUAGE UNDERSTANDING	7507
<i>Xuandi Fu, Carnegie Mellon University, United States of America; Feng-Ju Chang, Martin Radfar, Kai Wei, Jing Liu, Grant P. Strimel, Kanthashree Mysore Sathyendra, Amazon, United States of America</i>	
SPE-52.6: A BERT BASED JOINT LEARNING MODEL WITH FEATURE GATED MECHANISM FOR SPOKEN LANGUAGE UNDERSTANDING	7512
<i>Wang Zhang, Lei Jiang, Shaokang Zhang, Shuo Wang, Jianlong Tan, Chinese Academic of Science, China</i>	
SPE-53: SPEAKER RECOGNITION VII: NEURAL NETWORK ARCHITECTURE	
SPE-53.1: MFA: TDNN WITH MULTI-SCALE FREQUENCY-CHANNEL ATTENTION FOR TEXT-INDEPENDENT SPEAKER VERIFICATION WITH SHORT UTTERANCES	7517
<i>Tianchi Liu, Kong Aik Lee, Agency for Science, Technology and Research, Singapore; Rohan Kumar Das, Fortemedia Singapore, Singapore; Haizhou Li, National University of Singapore, Singapore</i>	
SPE-53.2: MLP-SVNET : A MULTI-LAYER PERCEPTRONS BASED NETWORK FOR SPEAKER VERIFICATION	7522
<i>Bing Han, Zhengyang Chen, Bei Liu, Yanmin Qian, Shanghai Jiao Tong University, China</i>	
SPE-53.3: REAL ADDITIVE MARGIN SOFTMAX FOR SPEAKER VERIFICATION	7527
<i>Lantian Li, Ruiqian Nai, Dong Wang, Tsinghua University, China</i>	
SPE-53.4: STATISTICAL PYRAMID DENSE TIME DELAY NEURAL NETWORK FOR SPEAKER VERIFICATION	7532
<i>Zi-Kai Wan, Qing-Hua Ren, You-Cai Qin, Qi-Rong Mao, Jiangsu University, China</i>	
SPE-53.5: ON THE IMPORTANCE OF DIFFERENT FREQUENCY BINS FOR SPEAKER VERIFICATION	7537
<i>Aiwen Deng, Wenxiong Kang, Feiqi Deng, South China University of Technology, China; Shuai Wang, Shanghai Jiao Tong University, China</i>	
SPE-53.6: SELF-KNOWLEDGE DISTILLATION VIA FEATURE ENHANCEMENT FOR SPEAKER VERIFICATION	7542
<i>Bei Liu, Haoyu Wang, Zhengyang Chen, Shuai Wang, Yanmin Qian, Shanghai Jiao Tong University, China</i>	
SPE-54: KEYWORD SPOTTING	
SPE-54.1: JOINT EGO-NOISE SUPPRESSION AND KEYWORD SPOTTING ON SWEEPING ROBOTS	7547
<i>Yueyue Na, Ziteng Wang, Qiang Fu, Alibaba Group, China; Liang Wang, Sun Yat-sen University (SYSU), China</i>	
SPE-54.2: PROGRESSIVE CONTINUAL LEARNING FOR SPOKEN KEYWORD SPOTTING	7552
<i>Yizheng Huang, Nancy Chen, Institute for Infocomm Research, A*STAR, Singapore, Singapore; Nana Hou, Nanyang Technological University, Singapore, Singapore</i>	

SPE-54.3: UNIFIED SPECULATION, DETECTION, AND VERIFICATION KEYWORD SPOTTING	7557
<i>Geng-shen Fu, Thibaud Senechal, Aaron Challenner, Tao Zhang, Amazon, United States of America</i>	
SPE-54.4: LEARNING DECOUPLING FEATURES THROUGH ORTHOGONALITY REGULARIZATION	7562
<i>Li Wang, Rongzhi Gu, Yuexian Zou, Peking University, China; Weiji Zhuang, Peng Gao, Yujun Wang, Xiaomi Inc., China</i>	
SPE-54.5: TEMPORAL EARLY EXITING FOR STREAMING SPEECH COMMANDS RECOGNITION	7567
<i>Raphael Tang, Karun Kumar, Piyush Vyas, Gefei Yang, Yajie Mao, Craig Murray, Comcast, United States of America; Ji Xin, Jimmy Lin, University of Waterloo, Canada; Wenyan Li, SenseTime, China</i>	
SPE-54.6: A STUDY OF DESIGNING COMPACT AUDIO-VISUAL WAKE WORD SPOTTING SYSTEM BASED ON ITERATIVE FINE-TUNING IN NEURAL NETWORK PRUNING	7572
<i>Hengshun Zhou, Jun Du, University of Science and Technology of China, China; Chao-Han Huck Yang, Chin-Hui Lee, Georgia Institute of Technology, United States of America; Shifu Xiong, iFlytek Research, China</i>	
SPE-55: SPEECH SYNTHESIS: PROSODY	
SPE-55.1: PROSOSPEECH: ENHANCING PROSODY WITH QUANTIZED VECTOR PRE-TRAINING IN TEXT-TO-SPEECH	7577
<i>Yi Ren, Zhou Zhao, Zhejiang University, China; Ming Lei, Zhiying Huang, Shiliang Zhang, Qian Chen, Zhijie Yan, Alibaba Group, China</i>	
SPE-55.2: PROSODY SPEECH: TOWARDS ADVANCED PROSODY MODEL FOR NEURAL TEXT-TO-SPEECH	7582
<i>Yuanhao Yi, Lei He, Shifeng Pan, Xi Wang, Yujia Xiao, Microsoft, China</i>	
SPE-55.3: HIERARCHICAL PROSODY MODELING AND CONTROL IN NON-AUTOREGRESSIVE PARALLEL NEURAL TTS	7587
<i>Tuomo Raitio, Jiangchuan Li, Shreyas Seshadri, Apple, United States of America</i>	
SPE-55.4: DISCOURSE-LEVEL PROSODY MODELING WITH A VARIATIONAL AUTOENCODER FOR NON-AUTOREGRESSIVE EXPRESSIVE SPEECH SYNTHESIS	7592
<i>Ning-Qian Wu, Zhao-Ci Liu, Zhen-Hua Ling, University of Science and Technology of China, China</i>	
SPE-55.5: UNSUPERVISED WORD-LEVEL PROSODY TAGGING FOR CONTROLLABLE SPEECH SYNTHESIS	7597
<i>Yiwei Guo, Chenpeng Du, Kai Yu, Shanghai Jiaotong University, China</i>	
SPE-55.6: A CHARACTER-LEVEL SPAN-BASED MODEL FOR MANDARIN PROSODIC STRUCTURE PREDICTION	7602
<i>Xueyuan Chen, Changhe Song, Yixuan Zhou, Zhiyong Wu, Tsinghua University, China; Changbin Chen, Zhongqin Wu, TAL Education Group, China; Helen Meng, The Chinese University of Hong Kong, China</i>	
SPE-56: LANGUAGE UNDERSTANDING III: INTENT DETECTION AND SLOT FILLING	
SPE-56.1: SLIM: EXPLICIT SLOT-INTENT MAPPING WITH BERT FOR JOINT MULTI-INTENT DETECTION AND SLOT FILLING	7607
<i>Fengyu Cai, Wanhao Zhou, Boi Faltings, EPFL, Switzerland; Fei Mi, Huawei Noah's Ark Lab, China</i>	
SPE-56.2: JOINT MULTIPLE INTENT DETECTION AND SLOT FILLING VIA SELF-DISTILLATION	7612
<i>Lisong Chen, Peilin Zhou, Yuexian Zou, Peking University, China</i>	

SPE-56.3: A GRAPH ATTENTION INTERACTIVE REFINE FRAMEWORK WITH	7617
CONTEXTUAL REGULARIZATION FOR JOINTING INTENT DETECTION AND SLOT FILLING	
<i>Zhanbiao Zhu, Peijie Huang, Haojing Huang, Shudong Liu, Leyi Lao, South China Agricultural University, China</i>	
SPE-56.4: ADJACENCY PAIRS-AWARE HIERARCHICAL ATTENTION NETWORKS FOR	7622
DIALOGUE INTENT CLASSIFICATION	
<i>Jiabao Xu, Peijie Huang, Youming Peng, Jiande Ding, Boxi Huang, Simin Huang, South China Agricultural University, China</i>	
SPE-56.5: ADVIN: AUTOMATICALLY DISCOVERING NOVEL DOMAINS AND INTENTS	7627
FROM USER TEXT UTTERANCES	
<i>Nikhita Vedula, Rahul Gupta, Aman Alok, Mukund Sridhar, Shankar Ananthakrishnan, Amazon, United States of America</i>	
SPE-56.6: A NEW DATA AUGMENTATION METHOD FOR INTENT CLASSIFICATION	7632
ENHANCEMENT AND ITS APPLICATION ON SPOKEN CONVERSATION DATASETS	
<i>Zvi Kons, Aharon Satt, Hong-Kwang Kuo, Samuel Thomas, Boaz Carmeli, Ron Hoory, Brian Kingsbury, IBM, Israel</i>	
 SPE-57: SPEAKER RECOGNITION VIII: ROBUSTNESS	
SPE-57.1: ROBUST SPEAKER VERIFICATION WITH JOINT SELF-SUPERVISED AND	7637
SUPERVISED LEARNING	
<i>Kai Wang, Xiaolei Zhang, Miao Zhang, Yuguang Li, Samsung R&D Institute China Xian, China; Jaeyun Lee, Kiho Cho, Sung-Un Park, Samsung Advanced Institute of Technology, Korea, Republic of</i>	
SPE-57.2: ROBUST SPEAKER VERIFICATION USING POPULATION-BASED DATA	7642
AUGMENTATION	
<i>Weiwei Lin, Man-Wai Mak, The Hong Kong Polytechnic University, Hong Kong</i>	
SPE-57.3: RAWNEXT: SPEAKER VERIFICATION SYSTEM FOR VARIABLE-DURATION	7647
UTTERANCES WITH DEEP LAYER AGGREGATION AND EXTENDED DYNAMIC SCALING	
POLICIES	
<i>Ju-ho Kim, Hye-jin Shim, Jungwoo Heo, Ha-Jin Yu, University of Seoul, Korea, Republic of</i>	
SPE-57.4: CONTRASTIVE-MIXUP LEARNING FOR IMPROVED SPEAKER	7652
VERIFICATION	
<i>Xin Zhang, Texas A&M University, United States of America; Minh Jin, Roger Cheng, Ruirui Li, Eunjung Han, Andreas Stolcke, Amazon Inc., United States of America</i>	
SPE-57.5: A STUDY OF THE ROBUSTNESS OF RAW WAVEFORM BASED SPEAKER	7657
EMBEDDINGS UNDER MISMATCHED CONDITIONS	
<i>Ge Zhu, Frank Cwitkowitz, Zhiyao Duan, University of Rochester, United States of America</i>	
SPE-57.6: DISENTANGLED SPEAKER EMBEDDING FOR ROBUST SPEAKER	7662
VERIFICATION	
<i>Lu Yi, Man-Wai Mak, The Hong Kong Polytechnic University, Hong Kong</i>	
 SPE-58: SPEECH RECOGNITION: ACOUSTIC MODELING III	
SPE-58.1: PERFORMANCE-EFFICIENCY TRADE-OFFS IN UNSUPERVISED	7667
PRE-TRAINING FOR SPEECH RECOGNITION	
<i>Felix Wu, Kwangyoun Kim, Jing Pan, Kyu Han, ASAPP, United States of America; Kilian Weinberger, Yoav Artzi, ASAPP and Cornell University, United States of America</i>	

SPE-58.2: ADVANCING MOMENTUM PSEUDO-LABELING WITH CONFORMER AND INITIALIZATION STRATEGY	7672
<i>Yosuke Higuchi, Waseda University, Mitsubishi Electric Research Laboratories, Japan; Niko Moritz, Facebook AI, United Kingdom of Great Britain and Northern Ireland; Jonathan Le Roux, Takaaki Hori, Mitsubishi Electric Research Laboratories, United States of America</i>	
SPE-58.3: TTS4PRETRAIN 2.0: ADVANCING THE USE OF TEXT AND SPEECH IN ASR PRETRAINING WITH CONSISTENCY AND CONTRASTIVE LOSSES	7677
<i>Zhehuai Chen, Yu Zhang, Andrew Rosenberg, Bhuvana Ramabhadran, Pedro Moreno, Gary Wang, Google, United States of America</i>	
SPE-58.4: SYNT + + : UTILIZING IMPERFECT SYNTHETIC DATA TO IMPROVE SPEECH RECOGNITION	7682
<i>Ting-Yao Hu, Ashish Shrivastava, Jen-Hao Rick Chang, Hema Koppula, Oncel Tuzel, Apple, United States of America; Mohammadreza Armandpour, Texas A&M University, United States of America</i>	
SPE-58.5: PSEUDO-LABELING FOR MASSIVELY MULTILINGUAL SPEECH RECOGNITION	7687
<i>Loren Lugosch, McGill University/Mila, Canada; Tatiana Likhomanenko, Gabriel Synnaeve, Ronan Collobert, Facebook AI Research, United States of America</i>	
SPE-58.6: DP-DWA: DUAL-PATH DYNAMIC WEIGHT ATTENTION NETWORK WITH STREAMING DFSMN-SAN FOR AUTOMATIC SPEECH RECOGNITION	7692
<i>Dongpeng Ma, Yiwen Wang, Liqiang He, Mingjie Jin, Dan Su, Dong Yu, Tencent, China</i>	
SPE-59: EMOTION RECOGNITION: GENERAL TOPICS II	
SPE-59.2: SERAB: A MULTI-LINGUAL BENCHMARK FOR SPEECH EMOTION RECOGNITION	7697
<i>Neil Scheidwasser-Clow, Pierre Beckmann, École Polytechnique Fédérale de Lausanne (EPFL), Switzerland; Mikolaj Kegler, Imperial College London, United Kingdom of Great Britain and Northern Ireland; Milos Cernak, Logitech Europe S.A., Switzerland</i>	
SPE-59.3: ENHANCING PRIVACY THROUGH DOMAIN ADAPTIVE NOISE INJECTION FOR SPEECH EMOTION RECOGNITION	7702
<i>Tiantian Feng, Hanieh Hashemi, Murali Annavaram, Shrikanth Narayanan, University of Southern California, United States of America</i>	
SPE-59.4: SELECTIVE MULTI-TASK LEARNING FOR SPEECH EMOTION RECOGNITION USING CORPORA OF DIFFERENT STYLES	7707
<i>Heran Zhang, Masato Mimura, Tatsuya Kawahara, Kyoto University, Japan; Kenkichi Ishizuka, Revcomm, Japan</i>	
SPE-59.5: FRONTEND ATTRIBUTES DISENTANGLEMENT FOR SPEECH EMOTION RECOGNITION	7712
<i>Yu-Xuan Xi, Yan Song, Li-Rong Dai, National Engineering Laboratory for Speech and Language Information Processing, University of Science and Technology of China, China; Ian McLoughlin, ICT Cluster, Singapore Institute of Technology, Singapore; Lin Liu, iFLYTEK Research, iFLYTEK CO., LTD, China</i>	
SPE-59.6: EXPLOITING ANNOTATORS' TYPED DESCRIPTION OF EMOTION PERCEPTION TO MAXIMIZE UTILIZATION OF RATINGS FOR SPEECH EMOTION RECOGNITION	7717
<i>Huang-Cheng Chou, Wei-Cheng Lin, Carlos Busso, The University of Texas at Dallas, United States of America; Chi-Chun Lee, National Tsing Hua University,, Taiwan</i>	
SPE-60: MULTIMODAL LANGUAGE PROCESSING	
SPE-60.1: AUTOMATED AUDIO CAPTIONING USING TRANSFER LEARNING AND RECONSTRUCTION LATENT SPACE SIMILARITY REGULARIZATION	7722
<i>Andrew Koh, Fuzhao Xue, Eng Siong Chng, Nanyang Technological University, Singapore</i>	

SPE-60.2: FAST-SLOW TRANSFORMER FOR VISUALLY GROUNDING SPEECH	7727
<i>Puyuan Peng, David Harwath, The University of Texas at Austin, United States of America</i>	
SPE-60.3: AUDIO-VISUAL SCENE-AWARE DIALOG AND REASONING USING	7732
AUDIO-VISUAL TRANSFORMERS WITH JOINT STUDENT-TEACHER LEARNING	
<i>Ankit Parag Shah, Carnegie Mellon University, United States of America; Shijie Geng, Rutgers University, United States of America; Gao Peng, Chinese University of Hong Kong, United States of America; Anoop Cherian, Takaaki Hori, Tim K. Marks, Jonathan Le Roux, Chiori Hori, Mitsubishi Electric Research Laboratories (MERL), United States of America</i>	
SPE-60.4: AIMNET: ADAPTIVE IMAGE-TAG MERGING NETWORK FOR AUTOMATIC	7737
MEDICAL REPORT GENERATION	
<i>Jijun Shi, Shanshe Wang, Ronggang Wang, Siwei Ma, Peking University, China</i>	
SPE-60.5: ADVERSARIAL INPUT ABLATION FOR AUDIO-VISUAL LEARNING	7742
<i>David Xu, David Harwath, The University of Texas at Austin, United States of America</i>	
SPE-60.6: GATED MULTIMODAL FUSION WITH CONTRASTIVE LEARNING FOR	7747
TURN-TAKING PREDICTION IN HUMAN-ROBOT DIALOGUE	
<i>Jiudong Yang, Peiying Wang, Mingchao Feng, Meng Chen, Xiaodong He, JD AI, China; Yi Zhu, University of Cambridge, United Kingdom of Great Britain and Northern Ireland</i>	
 SPE-61: SPEECH ENHANCEMENT AND CODING	
SPE-61.1: JOINT FAR- AND NEAR-END SPEECH INTELLIGIBILITY ENHANCEMENT	7752
BASED ON THE APPROXIMATED SPEECH INTELLIGIBILITY INDEX	
<i>Andreas Jonas Fuglsig, Lars Søndergaard Bertelsen, Peter Mariager, RTX A/S, Denmark; Jan Østergaard, Jesper Jensen, Zheng-Hua Tan, Aalborg University, Denmark</i>	
SPE-61.2: ATTENTION-BASED FUSION FOR BONE-CONDUCTED AND	7757
AIR-CONDUCTED SPEECH ENHANCEMENT IN THE COMPLEX DOMAIN	
<i>Heming Wang, DeLiang Wang, The Ohio State University, United States of America; Xueliang Zhang, Inner Mongolia University, China</i>	
SPE-61.3: A TWO-STEP BACKWARD COMPATIBLE FULLBAND SPEECH	7762
ENHANCEMENT SYSTEM	
<i>Xu Zhang, Lianwu Chen, Xiguang Zheng, Xinlei Ren, Chen Zhang, Liang Guo, Bing Yu, Kuai Shou Technology Co. Beijing, China, China</i>	
SPE-61.4: S-DCCRN: SUPER WIDE BAND DCCRN WITH LEARNABLE COMPLEX	7767
FEATURE FOR SPEECH ENHANCEMENT	
<i>Shubo Lv, Yihui Fu, Mengtao Xing, Jiayao Sun, Lei Xie, Northwestern Polytechnical University, Xi'an, China, China; Jun Huang, Yannan Wang, Tao Yu, Tencent Ethereal Audio Lab, Tencent Corporation, Shenzhen, China, China</i>	
SPE-61.5: COGNITIVE CODING OF SPEECH	7772
<i>Reza Lotfidereshgi, Philippe Gournay, Universite de Sherbrooke, Canada</i>	
SPE-61.6: SPEECH ENHANCEMENT FOR LOW BIT RATE SPEECH CODEC	7777
<i>Ju Lin, Clemson University, United States of America; Kaustubh Kalgaonkar, Qing He, Xin Lei, Facebook, United States of America</i>	

SPE-62: SPEECH RECOGNITION: TRAINING METHODS FOR E2E ASR

SPE-62.1: CONSISTENT TRAINING AND DECODING FOR END-TO-END SPEECH	7782
RECOGNITION USING LATTICE-FREE MMI	
<i>Jinchuan Tian, Yuexian Zou, Peking University, China; Jianwei Yu, Chao Weng, Shi-Xiong Zhang, Dan Su, Dong Yu, Tencent, China</i>	

SPE-62.2: BEING GREEDY DOES NOT HURT: SAMPLING STRATEGIES FOR END-TO-END SPEECH RECOGNITION	7787
<i>Jahn Heymann, Egor Lakomkin, Leif Raedel, Amazon Inc., Germany</i>	
SPE-62.3: INVESTIGATING SEQUENCE-LEVEL NORMALISATION FOR CTC-LIKE END-TO-END ASR	7792
<i>Zeyu Zhao, Peter Bell, Centre for Speech Technology Research, University of Edinburgh, UK, United Kingdom of Great Britain and Northern Ireland</i>	
SPE-62.4: HIERARCHICAL CONDITIONAL END-TO-END ASR WITH CTC AND MULTI-GRANULAR SUBWORD UNITS	7797
<i>Yosuke Higuchi, Keita Karube, Tetsuji Ogawa, Tetsunori Kobayashi, Waseda University, Japan</i>	
SPE-62.5: OPTIMIZING ALIGNMENT OF SPEECH AND LANGUAGE LATENT SPACES FOR END-TO-END SPEECH RECOGNITION AND UNDERSTANDING	7802
<i>Wei Wang, Yanmin Qian, Shanghai Jiao Tong University, China; Shuo Ren, Yao Qian, Shujie Liu, Yu Shi, Michael Zeng, Microsoft Corporation, China</i>	
SPE-62.6: MINIMUM WORD ERROR TRAINING FOR NON-AUTOREGRESSIVE TRANSFORMER-BASED CODE-SWITCHING ASR	7807
<i>Yizhou Peng, Jicheng Zhang, Hao Huang, Xinjiang University, China; Haihua Xu, Bytedance, Singapore; Eng Siong Chng, Nanyang Technological University, Singapore</i>	
SPE-63: TEXT CLASSIFICATION	
SPE-63.1: CONNECTING TARGETS VIA LATENT TOPICS AND CONTRASTIVE LEARNING: A UNIFIED FRAMEWORK FOR ROBUST ZERO-SHOT AND FEW-SHOT STANCE DETECTION	7812
<i>Rui Liu, Zheng Lin, Peng Fu, Yuanxin Liu, Weiping Wang, Institute of Information Engineering, Chinese Academy of Sciences, China</i>	
SPE-63.2: PRIOR-BERT AND MULTI-TASK LEARNING FOR TARGET-ASPECT-SENTIMENT JOINT DETECTION	7817
<i>Cai Ke, Qingyu Xiong, Chao Wu, Hualing Yi, Chongqing University, China; Zikai Liao, Xidian University, China</i>	
SPE-63.3: CROSS-TARGET STANCE DETECTION VIA REFINED META-LEARNING	7822
<i>Huishan Ji, Zheng Lin, Peng Fu, Weiping Wang, Institute of Information Engineering, Chinese Academy of Sciences, China</i>	
SPE-63.4: A ROBUST CONTRASTIVE ALIGNMENT METHOD FOR MULTI-DOMAIN TEXT CLASSIFICATION	7827
<i>Xuefeng Li, Hao Lei, Liwen Wang, Guanting Dong, Jiachi Liu, Weiran Xu, Beijing University of Posts and Telecommunications, China; Jinzheng Zhao, University of Surrey, United Kingdom of Great Britain and Northern Ireland; Chunyun Zhang, Shandong University of Finance and Economics, China</i>	
SPE-63.5: INCREMENTAL USER EMBEDDING MODELING FOR PERSONALIZED TEXT CLASSIFICATION	7832
<i>Ruixue Lian, University of Wisconsin-Madison, United States of America; Che-Wei Huang, Yuqing Tang, Qilong Gu, Chengyuan Ma, Chenlei Guo, Amazon Alexa, United States of America</i>	
SPE-63.6: BLOCK-SPARSE ADVERSARIAL ATTACK TO FOOL TRANSFORMER-BASED TEXT CLASSIFIERS	7837
<i>Sahar Sadrizadeh, Pascal Frossard, École polytechnique fédérale de Lausanne (EPFL), Switzerland; Ljiljana Dolamic, Armasuisse S + T, Switzerland</i>	

SPE-64: SPEECH ENHANCEMENT: DNN ARCHITECTURES

SPE-64.1: MANNER: MULTI-VIEW ATTENTION NETWORK FOR NOISE ERASURE 7842
Hyun Joon Park, Byung Ha Kang, Wooseok Shin, Jin Sob Kim, Sung Won Han, Korea University, Korea, Republic of

SPE-64.2: DUAL-BRANCH ATTENTION-IN-ATTENTION TRANSFORMER FOR 7847
SINGLE-CHANNEL SPEECH ENHANCEMENT
Guochen Yu, Communication University of China\ Institute of Acoustics, Chinese Academy of Sciences, China; Andong Li, Chengshi Zheng, Institute of Acoustics, Chinese Academy of Sciences, China; YINUO Guo, Bytedance, China; Yutian Wang, Hui Wang, Communication University of China, China

SPE-64.3: TIME-FREQUENCY ATTENTION FOR MONAURAL SPEECH 7852
ENHANCEMENT
Qiquan Zhang, Haizhou Li, Department of Electrical and Computer Engineering, National University of Singapore, Singapore, Singapore; Qi Song, Alibaba Group, China, China; Zhaozheng Ni, Meta AI, United States of America; Aaron Nicolson, Australian eHealth Research Centre, CSIRO, Herston, QLD, 4006, Australia, Australia

SPE-64.4: FULLSUBNET + : CHANNEL ATTENTION FULLSUBNET WITH COMPLEX 7857
SPECTROGRAMS FOR SPEECH ENHANCEMENT
Jun Chen, Zilin Wang, Zhiyong Wu, Tsinghua University, China; Deyi Tuo, Shiyin Kang, Huya Inc., China; Helen Meng, The Chinese University of Hong Kong, China

SPE-64.5: CROSS-DOMAIN SPEECH ENHANCEMENT WITH A NEURAL CASCADE 7862
ARCHITECTURE
Heming Wang, DeLiang Wang, The Ohio State University, United States of America

SPE-64.6: SPEECH DENOISING IN THE WAVEFORM DOMAIN WITH 7867
SELF-ATTENTION
Zhifeng Kong, University of California San Diego, United States of America; Wei Ping, Ambrish Dantrey, Bryan Catanzaro, Nvidia, United States of America

SPE-65: SPEECH RECOGNITION: GENERAL TOPICS IV

SPE-65.1: SRU + + : PIONEERING FAST RECURRENCE WITH ATTENTION FOR 7872
SPEECH RECOGNITION
Jing Pan, Tao Lei, Kwangyoun Kim, Kyu Han, ASAPP Inc, United States of America; Shinji Watanabe, Carnegie Mellon University, United States of America

SPE-65.2: LATTENTION: LATTICE-ATTENTION IN ASR RESCORING..... 7877
Prabhat Pandey, Sergio Duarte Torres, Ali Orkan Bayer, Ankur Gandhe, Volker Leutnant, Amazon, Germany

SPE-65.4: LEARNING ACOUSTIC FRAME LABELING FOR PHONEME 7882
SEGMENTATION WITH REGULARIZED ATTENTION MECHANISM
Binghuai Lin, Liyuan Wang, Tencent Technology Co., Ltd, China

SPE-65.5: LISTEN, KNOW AND SPELL: KNOWLEDGE-INFUSED SUBWORD 7887
MODELING FOR IMPROVING ASR PERFORMANCE OF OOV NAMED ENTITIES
Nilaksh Das, Duen Horng Chau, Georgia Institute of Technology, United States of America; Monica Sunkara, Dhanush Bekal, Sravan Bodapati, Katrin Kirchhoff, Amazon AWS AI, United States of America

SPE-65.6: JOINT SPEECH RECOGNITION AND AUDIO CAPTIONING 7892
Chaitanya Narisetty, Xuankai Chang, Shinji Watanabe, Carnegie Mellon University, United States of America; Emiru Tsunoo, Yosuke Kashiwagi, Michael Hentschel, Sony Group Corporation, Japan

SPE-66: SPEECH SYNTHESIS: STYLE CONTROL, TRANSFER, AND ADAPTATION

SPE-66.1: SPEAKER GENERATION 7897
Daisy Stanton, Matt Shannon, Soroosh Mariooryad, RJ Skerry-Ryan, Eric Battenberg, Tom Bagby, David Kao, Google, United States of America

SPE-66.2: VOICE FILTER: FEW-SHOT TEXT-TO-SPEECH SPEAKER ADAPTATION 7902
USING VOICE CONVERSION AS A POST-PROCESSING MODULE
Adam Gabryś, Goeric Huybrechts, Manuel Sam Ribeiro, Julian Roth, Giulia Comini, Roberto Barra-Chicote, Bartek Perz, Jaime Lorenzo-Trueba, Amazon, Poland; Chung-Ming Chien, National Taiwan University (NTU), United Kingdom of Great Britain and Northern Ireland

SPE-66.3: FINE-GRAINED STYLE CONTROL IN TRANSFORMER-BASED 7907
TEXT-TO-SPEECH SYNTHESIS
Li-Wei Chen, Alexander Rudnicky, Carnegie Mellon University, United States of America

SPE-66.4: USING MULTIPLE REFERENCE AUDIOS AND STYLE EMBEDDING 7912
CONSTRAINTS FOR SPEECH SYNTHESIS
Cheng Gong, Longbiao Wang, Tianjin University, China; Zhenhua Ling, University of Science and Technology of China, China; Ju Zhang, Huiyan Technology (Tianjin) Co., Ltd, China; Jianwu Dang, Japan Advanced Institute of Science and Technology, Japan

SPE-66.5: ENHANCING SPEAKING STYLES IN CONVERSATIONAL 7917
TEXT-TO-SPEECH SYNTHESIS WITH GRAPH-BASED MULTI-MODAL CONTEXT MODELING
Jingbei Li, Yi Meng, Chenyi Li, Zhiyong Wu, Tsinghua University, China; Helen Meng, The Chinese University of Hong Kong, China; Chao Weng, Dan Su, Tencent, China

SPE-66.6: TOWARDS EXPRESSIVE SPEAKING STYLE MODELLING WITH 7922
HIERARCHICAL CONTEXT INFORMATION FOR MANDARIN SPEECH SYNTHESIS
Shun Lei, Yixuan Zhou, Liyang Chen, Zhiyong Wu, Tsinghua University, China; Shiyin Kang, Huya Inc., China; Helen Meng, The Chinese University of Hong Kong, Hong Kong

SPE-67: LANGUAGE UNDERSTANDING IV: GENERAL TOPICS

SPE-67.1: SLUE: NEW BENCHMARK TASKS FOR SPOKEN LANGUAGE 7927
UNDERSTANDING EVALUATION ON NATURAL SPEECH
Suwon Shon, Felix Wu, Pablo Brusco, Kyu J. Han, ASAPP, United States of America; Ankita Pasad, Karen Livescu, Toyota Technological Institute at Chicago, United States of America; Yoav Artzi, Cornell University, ASAPP, United States of America

SPE-67.2: TOWARDS REDUCING THE NEED FOR SPEECH TRAINING DATA TO 7932
BUILD SPOKEN LANGUAGE UNDERSTANDING SYSTEMS
Samuel Thomas, Hong-Kwang Kuo, Brian Kingsbury, George Saon, IBM Research AI, United States of America

SPE-67.3: IMPROVING CROSS-MODAL UNDERSTANDING IN VISUAL DIALOG VIA 7937
CONTRASTIVE LEARNING
Feilong Chen, Xiuyi Chen, Shuang Xu, Bo Xu, Institute of Automation, Chinese Academy of Sciences (CASIA), China

SPE-67.4: NEWS RECOMMENDATION VIA MULTI-INTEREST NEWS SEQUENCE 7942
MODELLING
Rongyao Wang, Wenpeng Lu, Qilu University of Technology (Shandong Academy of Sciences), China; Shoujin Wang, RMIT University, Australia; Xueping Peng, University of Technology Sydney, Australia

SPE-67.5: MULTI-LEVEL CONTRASTIVE LEARNING FOR CROSS-LINGUAL 7947
ALIGNMENT
Beiduo Chen, Wu Guo, Bin Gu, University of Science and Technology of China, China; Quan Liu, Yongchao Wang, iFLYTEK Research, China

SPE-67.6: AUGMENTATION STRATEGY OPTIMIZATION FOR LANGUAGE UNDERSTANDING	7952
<i>Chang-Ting Chu, Mahdin Rohmatillah, Jen-Tzung Chien, National Yang Ming Chiao Tung University, Taiwan; Ching-hsien Lee, Industrial Technology Research Institute, Taiwan</i>	
SPE-68: SPEAKER RECOGNITION IX: SINGLE AND MULTI CHANNEL	
SPE-68.1: MULTI-FEATURE INTEGRATION FOR SPEAKER EMBEDDING EXTRACTION	7957
<i>Sreekanth Sankala, Shaik Mohammad Rafi B, Sri Rama Murty K, Indian Institute of Technology Hyderabad, India</i>	
SPE-68.2: LEARNABLE NONLINEAR COMPRESSION FOR ROBUST SPEAKER VERIFICATION	7962
<i>Xuechen Liu, University of Eastern Finland & Inria, Finland; Md Sahidullah, Inria, France; Tomi Kinnunen, University of Eastern Finland, Finland</i>	
SPE-68.3: FINE-TUNING WAV2VEC2 FOR SPEAKER RECOGNITION	7967
<i>Nik Vaessen, David A. van Leeuwen, Radboud University, Netherlands</i>	
SPE-68.4: GRAPH ATTENTIVE FEATURE AGGREGATION FOR TEXT-INDEPENDENT SPEAKER VERIFICATION	7972
<i>Hye-jin Shim, Jungwoo Heo, Ha-Jin Yu, University of Seoul, Korea, Republic of; Jae-han Park, Ga-Hui Lee, KT Corporation, Korea, Republic of</i>	
SPE-68.5: MULTISV: DATASET FOR FAR-FIELD MULTI-CHANNEL SPEAKER VERIFICATION	7977
<i>Ladislav Mošner, Oldřich Plchot, Lukáš Burget, Jan Černocký, Faculty of Information Technology, Brno University of Technology, Czechia</i>	
SPE-68.6: MULTI-CHANNEL SPEAKER VERIFICATION WITH CONV-TASNET BASED BEAMFORMER	7982
<i>Ladislav Mošner, Oldřich Plchot, Lukáš Burget, Jan Černocký, Faculty of Information Technology, Brno University of Technology, Czechia</i>	
SPE-69: ATTENTION MECHANISM FOR SPEECH MODELS	
SPE-69.1: LETR: A LIGHTWEIGHT AND EFFICIENT TRANSFORMER FOR KEYWORD SPOTTING	7987
<i>Kevin Ding, Martin Zong, Jiakui Li, Baoxiang Li, SenseTime, China</i>	
SPE-69.2: COMPRESSING TRANSFORMER-BASED ASR MODEL BY TASK-DRIVEN LOSS AND ATTENTION-BASED MULTI-LEVEL FEATURE DISTILLATION	7992
<i>Yongjie Lv, Longbiao Wang, Meng Ge, Kiyoshi Honda, Tianjin University, China; Sheng Li, Chenchen Ding, National Institute of Information & Communications Technology (NICT), Japan; Lixin Pan, Yuguang Wang, Huiyan Technology (TianJin) Co., LTD., China; Jianwu Dang, Japan Advanced Institute of Science and Technology, Japan</i>	
SPE-69.3: SPATIAL PROCESSING FRONT-END FOR DISTANT ASR EXPLOITING SELF-ATTENTION CHANNEL COMBINATOR	7997
<i>Dushyant Sharma, Rong Gong, James Fosburgh, Stanislav Kruchinin, Patrick Naylor, Ljubomir Milanovic, Nuance Communications Inc, United States of America</i>	
SPE-69.4: EFFICIENT SEQUENCE TRAINING OF ATTENTION MODELS USING APPROXIMATIVE RECOMBINATION	8002
<i>Nils-Philipp Wynnands, Wilfried Michel, Jan Rosendahl, Ralf Schlüter, Hermann Ney, RWTH Aachen University, Germany</i>	

SPE-69.5: NEUFA: NEURAL NETWORK BASED END-TO-END FORCED ALIGNMENT WITH BIDIRECTIONAL ATTENTION MECHANISM	8007
<i>Jingbei Li, Yi Meng, Zhiyong Wu, Tsinghua University, China; Helen Meng, The Chinese University of Hong Kong, China; Qiao Tian, Yuping Wang, Yuxuan Wang, ByteDance, China</i>	
SPE-69.6: CONFORMER-BASED SPEECH RECOGNITION WITH LINEAR NYSTRÖM ATTENTION AND ROTARY POSITION EMBEDDING	8012
<i>Lahiru Samarakoon, Tsun-Yat Leung, Fano Labs, Hong Kong</i>	
SPE-70: SPEECH SYNTHESIS: MULTI-LINGUAL AND MULTIMODAL SYNTHESIS	
SPE-70.1: MULTILINGUAL TEXT-TO-SPEECH TRAINING USING CROSS LANGUAGE VOICE CONVERSION AND SELF-SUPERVISED LEARNING OF SPEECH REPRESENTATIONS	8017
<i>Jilong Wu, Adam Polyak, Yaniv Taigman, Prabhav Agrawal, Qing He, Facebook, United States of America; Jason Fong, The University of Edinburgh, United Kingdom of Great Britain and Northern Ireland</i>	
SPE-70.2: TOWARDS LIFELONG LEARNING OF MULTILINGUAL TEXT-TO-SPEECH SYNTHESIS	8022
<i>Mu Yang, University of Texas at Dallas, United States of America; Shaojin Ding, Texas A&M University, United States of America; Tianlong Chen, Zhangyang Wang, University of Texas at Austin, United States of America; Tong Wang, University of Science and Technology of China, China</i>	
SPE-70.3: ZERO-SHOT CROSS-LINGUAL TRANSFER USING MULTI-STREAM ENCODER AND EFFICIENT SPEAKER REPRESENTATION	8027
<i>Yibin Zheng, Zewang Zhang, Xinhui Li, Wenchao Su, Li Lu, Tencent Inc, China, China</i>	
SPE-70.4: VISUALTTS: TTS WITH ACCURATE LIP-SPEECH SYNCHRONIZATION FOR AUTOMATIC VOICE OVER	8032
<i>Junchen Lu, Mingyang Zhang, Haizhou Li, National University of Singapore, Singapore; Berrak Sisman, Singapore University of Technology and Design, Singapore; Rui Liu, National University of Singapore and Singapore University of Technology and Design, Singapore</i>	
SPE-70.5: DURATION MODELING OF NEURAL TTS FOR AUTOMATIC DUBBING	8037
<i>Johanes Effendi, Yogesh Virkar, Roberto Barra-Chicote, Marcello Federico, Amazon AI, Japan</i>	
SPE-70.6: LEARNING TO PREDICT SPEECH IN SILENT VIDEOS VIA AUDIOVISUAL ANALOGY	8042
<i>Ravindra Yadav, Vinay Namboodiri, Rajesh Hegde, Indian Institute of Technology Kanpur, India; Ashish Sardana, NVIDIA, India</i>	
SPE-71: TEXT COMPLETION AND SUMMARIZATION	
SPE-71.1: SELF-ATTENTION FOR INCOMPLETE UTTERANCE REWRITING	8047
<i>Yong Zhang, Zhitao Li, Jianzong Wang, Ning Cheng, Jing Xiao, Ping An Technology (Shenzhen) Co., Ltd., China</i>	
SPE-71.2: MULTI-TURN INCOMPLETE UTTERANCE RESTORATION AS OBJECT DETECTION	8052
<i>Wangjie Jiang, Siheng Li, Jiayi Li, Yujiu Yang, Tsinghua University, China</i>	
SPE-71.3: CLSEG: CONTRASTIVE LEARNING OF STORY ENDING GENERATION	8057
<i>Yuqiang Xie, Yue Hu, Luxi Xing, Yunpeng Li, Wei Peng, Ping Guo, Institute of Information Engineering, Chinese Academy of Sciences, China</i>	
SPE-71.4: EXPLICITLY MODELING IMPORTANCE AND COHERENCE FOR TIMELINE SUMMARIZATION	8062
<i>Qianren Mao, Jianxin Li, Jiazheng Wang, Xi Li, Peng Hao, Beihang University (BUAA), China; Lihong Wang, CNCERT/CC, China; Zheng Wang, University of Leeds, United Kingdom of Great Britain and Northern Ireland</i>	

SPE-71.5: TED TALK TEASER GENERATION WITH PRE-TRAINED MODELS	8067
<i>Gianluca Vico, Jan Niehues, Maastricht University, Netherlands</i>	
SPE-71.6: END-TO-END SPEECH SUMMARIZATION USING RESTRICTED SELF-ATTENTION	8072
<i>Roshan Sharma, Shruti Palaskar, Alan W Black, Florian Metze, Carnegie Mellon University, United States of America</i>	
SPE-72: SPEAKER DIARIZATION I	
SPE-72.1: TURN-TO-DIARIZE: ONLINE SPEAKER DIARIZATION CONSTRAINED BY TRANSFORMER TRANSDUCER SPEAKER TURN DETECTION	8077
<i>Wei Xia, University of Texas at Dallas, United States of America; Han Lu, Quan Wang, Anshuman Tripathi, Yiling Huang, Ignacio Lopez Moreno, Hasim Sak, Google, United States of America</i>	
SPE-72.2: TRANSCRIBE-TO-DIARIZE: NEURAL SPEAKER DIARIZATION FOR UNLIMITED NUMBER OF SPEAKERS USING END-TO-END SPEAKER-ATTRIBUTED ASR	8082
<i>Naoyuki Kanda, Xiong Xiao, Yashesh Gaur, Xiaofei Wang, Zhong Meng, Zhuo Chen, Takuya Yoshioka, Microsoft, United States of America</i>	
SPE-72.3: A MULTITASK LEARNING FRAMEWORK FOR SPEAKER CHANGE DETECTION WITH CONTENT INFORMATION FROM UNSUPERVISED SPEECH DECOMPOSITION	8087
<i>Hang Su, Danyang Zhao, Long Dang, Xixin Wu, Xunying Liu, Helen Meng, The Chinese University of Hong Kong, Hong Kong; Minglei Li, Huawei Cloud, China</i>	
SPE-72.4: ASR-AWARE END-TO-END NEURAL DIARIZATION.....	8092
<i>Aparna Khare, Eunjung Han, Yuguang Yang, Andreas Stolcke, Amazon, United States of America</i>	
SPE-72.5: REFORMULATING SPEAKER DIARIZATION AS COMMUNITY DETECTION WITH EMPHASIS ON TOPOLOGICAL STRUCTURE	8097
<i>Siqi Zheng, Hongbin Suo, Alibaba Group, China</i>	
SPE-72.6: TITANET: NEURAL MODEL FOR SPEAKER REPRESENTATION WITH 1D DEPTH-WISE SEPARABLE CONVOLUTIONS AND GLOBAL CONTEXT	8102
<i>Nithin Rao Koluguri, Taejin Park, Boris Ginsburg, NVIDIA, United States of America</i>	
SPE-73: SPEECH RECOGNITION: NEURAL TRANSDUCER MODELS	
SPE-73.1: TRANSDUCER-BASED STREAMING DELIBERATION FOR CASCADED ENCODERS	8107
<i>Ke Hu, Tara Sainath, Arun Narayanan, Ruoming Pang, Trevor Strohman, Google, United States of America</i>	
SPE-73.2: IMPROVING THE LATENCY AND QUALITY OF CASCADED ENCODERS.....	8112
<i>Tara Sainath, Yanzhang He, Arun Narayanan, Rami Botros, Weiran Wang, David Qiu, Chung-cheng Chiu, Rohit Prabhavalkar, Alexander Gruenstein, Anmol Gulati, Bo Li, David Rybach, Emmanuel Guzman, Ian McGraw, James Qin, Krzysztof Choromanski, Qiao Liang, Robert David, Ruoming Pang, Shuoyiin Chang, Trevor Strohman, W. Ronny Huang, Wei Han, Yonghui Wu, Yu Zhang, Google, Inc., United States of America</i>	
SPE-73.3: IMPROVING THE FUSION OF ACOUSTIC AND TEXT REPRESENTATIONS IN RNN-T	8117
<i>Chao Zhang, Bo Li, Zhiyun Lu, Tara Sainath, Shuo-Yiin Chang, Google LLC, United Kingdom of Great Britain and Northern Ireland</i>	
SPE-73.4: ADAPTIVE DISCOUNTING OF IMPLICIT LANGUAGE MODELS IN RNN-TRANSDUCERS	8122
<i>Vinit Unni, Preethi Jyothi, Sunita Sarawagi, Indian Institute of Technology Bombay, India; Shreya Khare, Ashish Mittal, Samarth Bharadwaj, IBM Research, India</i>	

SPE-73.5: INTEGRATING TEXT INPUTS FOR TRAINING AND ADAPTING RNN TRANSDUCER ASR MODELS	8127
<i>Samuel Thomas, Brian Kingsbury, George Saon, Hong-Kwang Kuo, IBM Research AI, United States of America</i>	
SPE-73.6: FACTORIZED NEURAL TRANSDUCER FOR EFFICIENT LANGUAGE MODEL ADAPTATION	8132
<i>Xie Chen, Zhong Meng, Sarangarajan Parthasarathy, Jinyu Li, Microsoft Corporation, United States of America</i>	
SPE-74: REPRESENTATIONS AND RELATIONSHIPS IN LANGUAGE	
SPE-74.1: INTEGRATING DEPENDENCY TREE INTO SELF-ATTENTION FOR SENTENCE REPRESENTATION	8137
<i>Junhua Ma, Yuxuan Liu, Shangbo Zhou, Chongqing University, China; Jiajun Li, Xue Li, The University of Queensland, Australia</i>	
SPE-74.2: METRICBERT: TEXT REPRESENTATION LEARNING VIA SELF-SUPERVISED TRIPLET TRAINING	8142
<i>Itzik Malkiel, Dvir Ginzburg, Oren Barkan, Avi Caciularu, Yoni Weill, Noam Koenigstein, microsoft, Israel</i>	
SPE-74.4: END-TO-END NEURAL COREFERENCE RESOLUTION REVISITED: A SIMPLE YET EFFECTIVE BASELINE	8147
<i>Tuan Lai, University of Illinois at Urbana-Champaign, United States of America; Trung Bui, Adobe, United States of America; Doo-Soon Kim, Roku, Inc., United States of America</i>	
SPE-74.5: LOCAL CONTEXT INTERACTION-AWARE GLYPH-VECTORS FOR CHINESE SEQUENCE TAGGING	8152
<i>JunYu Lu, PingJian Zhang, South China University of Technology, China</i>	
SPE-75: SPEECH ANALYSIS	
SPE-75.2: DEEP LEARNING FOR PROMINENCE DETECTION IN CHILDREN'S READ SPEECH	8157
<i>Mithilesh Vaidya, Kamini Sabu, Preeti Rao, Indian Institute of Technology, Bombay, India</i>	
SPE-75.3: TOWARDS A COMMON SPEECH ANALYSIS ENGINE	8162
<i>Hagai Aronowitz, Itai Gat, Edmilson Morais, Weizhong Zhu, Ron Hoory, IBM Research AI, Israel</i>	
SPE-75.4: PHONE-TO-AUDIO ALIGNMENT WITHOUT TEXT: A SEMI-SUPERVISED APPROACH	8167
<i>Jian Zhu, David Jurgens, University of Michigan, United States of America; Cong Zhang, Radboud University, Netherlands</i>	
SPE-75.5: ATTACHMENT RECOGNITION IN SCHOOL-AGE CHILDREN: A MULTIMODAL APPROACH BASED ON LANGUAGE AND PARALANGUAGE ANALYSIS	8172
<i>Huda Alsofyani, Alessandro Vinciarelli, University of Glasgow, United Kingdom of Great Britain and Northern Ireland</i>	
SPE-75.6: DETERMINING THE BEST ACOUSTIC FEATURES FOR SMOKER IDENTIFICATION	8177
<i>Zhizhong Ma, Yuanhang Qiu, Feng Hou, Ruili Wang, Massey University, New Zealand; Joanna Ting Wai Chu, Christopher Bullen, University of Auckland, New Zealand</i>	
SPE-76: RESOURCE CONSTRAINED SPEECH MODELING	
SPE-76.1: END-TO-END LOW RESOURCE KEYWORD SPOTTING THROUGH CHARACTER RECOGNITION AND BEAM-SEARCH RE-SCORING	8182
<i>Ephrem Tibebe Mekonnen, University of Trento, Italy; Alessio Brutti, Daniele Falavigna, Fondazione Bruno Kessler, Italy</i>	

SPE-76.2: CURRICULUM OPTIMIZATION FOR LOW-RESOURCE SPEECH RECOGNITION	8187
<i>Anastasia Kuznetsova, Francis Tyers, Indiana University Bloomington, United States of America; Anurag Kumar, Jennifer Drexler Fox, Rev.com, United States of America</i>	
SPE-76.3: EXPLORING EFFECTIVE DATA UTILIZATION FOR LOW-RESOURCE SPEECH RECOGNITION	8192
<i>Zhikai Zhou, Wei Wang, Wangyou Zhang, Yanmin Qian, Shanghai Jiao Tong University, China</i>	
SPE-76.4: OMNI-SPARSITY DNN: FAST SPARSITY OPTIMIZATION FOR ON-DEVICE STREAMING E2E ASR VIA SUPERNET	8197
<i>Haichuan Yang, Yuan Shanguan, Dilin Wang, Meng Li, Pierce Chuang, Xiaohui Zhang, Ganesh Venkatesh, Ozlem Kalinli, Vikas Chandra, Facebook, United States of America</i>	
SPE-76.5: ANALYZING THE ROBUSTNESS OF UNSUPERVISED SPEECH RECOGNITION	8202
<i>Guan-Ting Lin, Chan-Jan Hsu, Da-Rong Liu, Hung-Yi Lee, National Taiwan University, Taiwan; Yu Tsao, Academia Sinica, Taiwan</i>	
SPE-76.6: INTERPRETING INTERMEDIATE CONVOLUTIONAL LAYERS IN UNSUPERVISED ACOUSTIC WORD CLASSIFICATION	8207
<i>Gašper Beguš, Alan Zhou, University of California, Berkeley, United States of America</i>	
SPE-77: QUESTION ANSWERING AND INFERENCE	
SPE-77.1: CONTEXT MODELING WITH EVIDENCE FILTER FOR MULTIPLE CHOICE QUESTION ANSWERING	8212
<i>Sicheng Yu, Jing Jiang, Singapore Management University, Singapore; Hao Zhang, Nanyang Technological University, Singapore; Wei Jing, Alibaba Damo Academy, China</i>	
SPE-77.2: FROM SHALLOW TO DEEP: COMPOSITIONAL REASONING OVER GRAPHS FOR VISUAL QUESTION ANSWERING	8217
<i>Zihao Zhu, University of Chinese Academy of Sciences, China</i>	
SPE-77.3: A QUESTION-ORIENTED PROPAGATION NETWORK FOR NEWS READING COMPREHENSION	8222
<i>Liang Wen, Houfeng Wang, Yingwei Luo, Xiaolin Wang, Peking University, China; Dehong Ma, Jun Fan, Daiting Shi, Zhicong Cheng, Dawei Yin, Baidu Inc., China</i>	
SPE-77.4: SYNTAX-BASED GRAPH MATCHING FOR KNOWLEDGE BASE QUESTION ANSWERING	8227
<i>Lu Ma, Dan Luo, Xi Zhu, Institute of Information Engineering, Chinese Academy of Sciences; University of Chinese Academy of Sciences, China; Peng Zhang, Meilin Zhou, Qi Liang, Institute of Information Engineering, Chinese Academy of Sciences, China; Bin Wang, Xiaomi Inc., China</i>	
SPE-77.5: QA4QG: USING QUESTION ANSWERING TO CONSTRAIN MULTI-HOP QUESTION GENERATION	8232
<i>Dan Su, Peng Xu, Pascale Fung, Hong Kong University of Science and Technology, Hong Kong</i>	
SPE-77.6: PAIR-LEVEL SUPERVISED CONTRASTIVE LEARNING FOR NATURAL LANGUAGE INFERENCE	8237
<i>Shu'ang Li, Xuming Hu, Li Lin, Lijie Wen, Tsinghua University, China</i>	
SPE-78: SPEECH PRODUCTION AND PERCEPTION	
SPE-78.1: ACOUSTIC COMPARISON OF PHYSICAL VOCAL TRACT MODELS WITH HARD AND SOFT WALLS	8242
<i>Peter Birkholz, Patrick Häsner, Steffen Kürbis, TU Dresden, Germany</i>	

SPE-78.2: AN ERROR CORRECTION SCHEME FOR IMPROVED AIR-TISSUE BOUNDARY IN REAL-TIME MRI VIDEO FOR SPEECH PRODUCTION	8247
<i>Anwesha Roy, Varun Belagali, Prasanta Kumar Ghosh, Indian Institute of Science (IISc), Bangalore, India</i>	
SPE-78.3: REPEAT AFTER ME: SELF-SUPERVISED LEARNING OF ACOUSTIC-TO-ARTICULATORY MAPPING BY VOCAL IMITATION	8252
<i>Marc-Antoine Georges, Laurent Girin, Jean-Luc Schwartz, Thomas Hueber, GIPSA-lab, France; Julien Diard, LPNC, France</i>	
SPE-78.4: MULTI-SPEAKER PITCH TRACKING VIA EMBODIED SELF-SUPERVISED LEARNING	8257
<i>Xiang Li, Yifan Sun, Xihong Wu, Jing Chen, Peking University, China</i>	
SPE-78.5: IMPROVING THE CLASSIFICATION OF PHONETIC SEGMENTS FROM RAW ULTRASOUND USING SELF-SUPERVISED LEARNING AND HARD EXAMPLE MINING	8262
<i>Yunsheng Xiong, Kele Xu, Yong Dou, National University of Defense Technology, China; Meng Jiang, Changsha Medical University, China; Liang Cheng, Xiangya Hospital, Central South University, China; Jinjia Wang, Yanshan University, China</i>	
SPE-78.6: THE IMPACT OF CROSS LANGUAGE ON ACOUSTIC-TO-ARTICULATORY INVERSION AND ITS INFLUENCE ON ARTICULATORY SPEECH SYNTHESIS	8267
<i>Aravind Illa, Prasanta Kumar Ghosh, Indian Institute of Science, Bangalore, India; Aanish Nair, Georgia Institute of Technology, United States of America</i>	
SPE-79: SPEECH RECOGNITION: TRANSFORMERS	
SPE-79.1: TRANSFORMER-BASED STREAMING ASR WITH CUMULATIVE ATTENTION	8272
<i>Mohan Li, Shucong Zhang, Cătălin Zorilă, Rama Doddipatla, Toshiba Cambridge Research Laboratory, Toshiba Europe Ltd, United Kingdom of Great Britain and Northern Ireland</i>	
SPE-79.2: STREAMING TRANSFORMER TRANSDUCER BASED SPEECH RECOGNITION USING NON-CAUSAL CONVOLUTION	8277
<i>Yangyang Shi, Chunyang Wu, Dilin Wang, Alex Xiao, Jay Mahadeokar, Xiaohui Zhang, Chunxi Liu, Ke Li, Yuan Shangguan, Varun Nagaraja, Ozlem Kalinli, Mike Seltzer, Facebook AI, United States of America</i>	
SPE-79.3: HYBRID RNN-T/ATTENTION-BASED STREAMING ASR WITH TRIGGERED CHUNKWISE ATTENTION AND DUAL INTERNAL LANGUAGE MODEL INTEGRATION	8282
<i>Takafumi Moriya, Takanori Ashihara, Atsushi Ando, Hiroshi Sato, Tomohiro Tanaka, Kohei Matsuura, Ryo Masumura, Marc Delcroix, NTT Corporation, Japan; Takahiro Shinozaki, Tokyo Institute of Technology, Japan</i>	
SPE-79.4: RUN-AND-BACK STITCH SEARCH: NOVEL BLOCK SYNCHRONOUS DECODING FOR STREAMING ENCODER-DECODER ASR	8287
<i>Emiru Tsunoo, Michael Hentschel, Yosuke Kashiwagi, Sony Group Corporation, Japan; Chaitanya Narisetty, Shinji Watanabe, Carnegie Mellon University, United States of America</i>	
SPE-79.5: ALIGNMENT-LEARNING BASED SINGLE-STEP DECODING FOR ACCURATE AND FAST NON-AUTOREGRESSIVE SPEECH RECOGNITION	8292
<i>Yonghe Wang, Rui Liu, Feilong Bao, Hui Zhang, Guanglai Gao, Inner Mongolia University, China</i>	
SPE-79.6: USTED: IMPROVING ASR WITH A UNIFIED SPEECH AND TEXT ENCODER-DECODER	8297
<i>Bolaji Yusuf, Bogazici University, Turkey; Ankur Gandhe, Alex Sokolov, Amazon, United States of America</i>	
SPE-80: SPEECH SYNTHESIS: EXPRESSIVENESS AND VOICE CLONING	
SPE-80.1: IMPROVING EMOTIONAL SPEECH SYNTHESIS BY USING SUS-CONSTRAINED VAE AND TEXT ENCODER AGGREGATION	8302
<i>Fengyu Yang, Jian Luan, Yujun Wang, Xiaomi Corporation, China</i>	

SPE-80.2: DISTRIBUTION AUGMENTATION FOR LOW-RESOURCE EXPRESSIVE TEXT-TO-SPEECH	8307
<i>Mateusz Lajszczak, Animesh Prasad, Arent van Korlaar, Bajibabu Bollepalli, Antonio Bonafonte, Arnaud Joly, Marco Nicolis, Alexis Moinet, Thomas Drugman, Trevor Wood, Elena Sokolova, Amazon, United Kingdom of Great Britain and Northern Ireland</i>	
SPE-80.3: INTERACTIVE MULTI-LEVEL PROSODY CONTROL FOR EXPRESSIVE SPEECH SYNTHESIS	8312
<i>Tobias Cornille, Jessa Bekker, KU Leuven, Belgium; Fengna Wang, Acapela-Group, Belgium</i>	
SPE-80.4: IMPROVE FEW-SHOT VOICE CLONING USING MULTI-MODAL LEARNING	8317
<i>Haitong Zhang, Yue Lin, Netease Games AI Lab, China</i>	
SPE-80.5: CLONING ONE’S VOICE USING VERY LIMITED DATA IN THE WILD	8322
<i>Dongyang Dai, Yuanzhe Chen, Li Chen, Ming Tu, Lu Liu, Rui Xia, Qiao Tian, Yuping Wang, Yuxuan Wang, bytedance, China</i>	
SPE-80.6: UNET-TTS: IMPROVING UNSEEN SPEAKER AND STYLE TRANSFER IN ONE-SHOT VOICE CLONING	8327
<i>Rui Li, Dong Pu, Minnie Huang, Bill Huang, CloudMinds, China</i>	
SPE-81: ENTITY RECOGNITION	
SPE-81.1: IMPROVING BIOMEDICAL NAMED ENTITY RECOGNITION WITH A UNIFIED MULTI-TASK MRC FRAMEWORK	8332
<i>Yiqi Tong, Fuzhen Zhuang, Deqing Wang, Beihang University, China; Haochao Ying, Zhejiang University, China; Binling Wang, Xiamen University, China</i>	
SPE-81.2: A MULTI-TASK LEARNING FRAMEWORK FOR CHINESE MEDICAL PROCEDURE ENTITY NORMALIZATION	8337
<i>Xuhui Sui, Kehui Song, Baohang Zhou, Ying Zhang, Xiaojie Yuan, Nankai University, China</i>	
SPE-81.3: WASSERSTEIN CROSS-LINGUAL ALIGNMENT FOR NAMED ENTITY RECOGNITION	8342
<i>Rui Wang, Ricardo Henao, Duke University, United States of America</i>	
SPE-81.4: LEARNING COMMON DEPENDENCY STRUCTURE FOR UNSUPERVISED CROSS-DOMAIN NER	8347
<i>Luchen Liu, Bin Wang, Xiaomi Inc., China; Xixun Lin, Peng Zhang, Lei Zhang, Institute of Information Engineering, Chinese Academy of Sciences, China</i>	
SPE-81.5: AISHELL-NER: NAMED ENTITY RECOGNITION FROM CHINESE SPEECH	8352
<i>Boli Chen, Guangwei Xu, Xiaobin Wang, Pengjun Xie, Meishan Zhang, Fei Huang, Alibaba Group, China</i>	
SPE-81.6: CALL-SIGN RECOGNITION AND UNDERSTANDING FOR NOISY AIR-TRAFFIC TRANSCRIPTS USING SURVEILLANCE INFORMATION	8357
<i>Alexander Blatt, Dietrich Klakow, Saarland University, Germany; Martin Kocour, Karel Veselý, Igor Szöke, Brno University of Technology, Germany</i>	
SPE-82: SPEAKER DIARIZATION II	
SPE-82.1: INCORPORATING END-TO-END FRAMEWORK INTO TARGET-SPEAKER VOICE ACTIVITY DETECTION	8362
<i>Weiqing Wang, Duke University, United States of America; Ming Li, Duke Kunshan University, China</i>	

SPE-82.2: MULTI-SCALE SPEAKER EMBEDDING-BASED GRAPH ATTENTION NETWORKS FOR SPEAKER DIARISATION	8367
<i>Youngki Kwon, Hee-Soo Heo, Jee-weon Jung, You Jin Kim, Bong-Jin Lee, Naver Corporation, Korea, Republic of; Joon Son Chung, Korea Advanced Institute of Science and Technology, Korea, Republic of</i>	
SPE-82.3: TOWARDS END-TO-END SPEAKER DIARIZATION WITH GENERALIZED NEURAL SPEAKER CLUSTERING	8372
<i>Chunlei Zhang, Chao Weng, Meng Yu, Dong Yu, Tencent AI Lab, United States of America; Jiatong Shi, Carnegie Mellon University, United States of America</i>	
SPE-82.4: AUXILIARY LOSS OF TRANSFORMER WITH RESIDUAL CONNECTION FOR END-TO-END SPEAKER DIARIZATION	8377
<i>Yechan Yu, Dongkeon Park, Hong Kook Kim, Gwangju Institute of Science and Technology, Korea, Republic of</i>	
SPE-82.5: TIGHT INTEGRATION OF NEURAL- AND CLUSTERING-BASED DIARIZATION THROUGH DEEP UNFOLDING OF INFINITE GAUSSIAN MIXTURE MODEL	8382
<i>Keisuke Kinoshita, Marc Delcroix, Tomoharu Iwata, NTT, Japan</i>	
SPE-82.6: IMPROVING SEPARATION-BASED SPEAKER DIARIZATION VIA ITERATIVE MODEL REFINEMENT AND SPEAKER EMBEDDING BASED POST-PROCESSING	8387
<i>Shu-Tong Niu, Jun Du, University of Science and Technology of China, China; Lei Sun, iFlytek Research, China; Chin-Hui Lee, Georgia Institute of Technology, United States of America</i>	
SPE-83: NEW ALGORITHM FOR SPEECH RECOGNITION	
SPE-83.1: KNOWLEDGE DISTILLATION FROM LANGUAGE MODEL TO ACOUSTIC MODEL: A HIERARCHICAL MULTI-TASK LEARNING APPROACH	8392
<i>Munhak Lee, Joonhyuk Chang, Hanyang University, Korea, Republic of</i>	
SPE-83.2: IMPROVING PSEUDO-LABEL TRAINING FOR END-TO-END SPEECH RECOGNITION USING GRADIENT MASK	8397
<i>Shaoshi Ling, Chen Shen, Meng Cai, Zejun Ma, bytedance, China</i>	
SPE-83.3: MULTI-TURN RNN-T FOR STREAMING RECOGNITION OF MULTI-PARTY SPEECH	8402
<i>Ilya Sklyar, Anna Piunova, Amazon, Germany; Xianrui Zheng, University of Cambridge, United Kingdom of Great Britain and Northern Ireland; Yulan Liu, DeepMind, United Kingdom of Great Britain and Northern Ireland</i>	
SPE-83.4: ON LANGUAGE MODEL INTEGRATION FOR RNN TRANSDUCER BASED SPEECH RECOGNITION	8407
<i>Wei Zhou, Zuoyun Zheng, Ralf Schlüter, Hermann Ney, RWTH Aachen University, Germany</i>	
SPE-83.5: CACHING NETWORKS: CAPITALIZING ON COMMON SPEECH FOR ASR	8412
<i>Anastasios Alexandridis, Grant Strimel, Ariya Rastrow, Pavel Kveton, Jon Webb, Maurizio Omologo, Siegfried Kunzmann, Athanasios Mouchtaris, Amazon.com, United States of America</i>	
SPE-83.6: GPU-ACCELERATED FORWARD-BACKWARD ALGORITHM WITH APPLICATION TO LATTICE-FREE MMI	8417
<i>Lucas Ondel, Léa-Marie Lam-Yee-Mui, Caio Corro, Laboratoire Interdisciplinaire des Sciences du Numérique, France; Martin Kocour, Lukas Burget, Brno University of Technology, France</i>	
SPE-84: SPEECH SYNTHESIS: VOCODER AND EVALUATION	
SPE-84.1: ITOWAVE: ITO STOCHASTIC DIFFERENTIAL EQUATION IS ALL YOU NEED FOR WAVE GENERATION	8422
<i>Shoule Wu, Yangzhou University, China; Ziqiang Shi, Fujitsu R & D Center, China</i>	

SPE-84.2: MULTI-SAMPLE SUBBAND WAVERNN VIA MULTIVARIATE GAUSSIAN	8427
<i>Hiroki Kanagawa, Yusuke Ijima, NTT Corporation, Japan</i>	
SPE-84.3: INFERGRAD: IMPROVING DIFFUSION MODELS FOR VOCODER BY CONSIDERING INFERENCE IN TRAINING	8432
<i>Zehua Chen, Danilo Mandic, Imperial College London, United Kingdom of Great Britain and Northern Ireland; Xu Tan, Ke Wang, Shifeng Pan, Lei He, Sheng Zhao, Microsoft, China; , ,</i>	
SPE-84.4: NEURAL SPEECH SYNTHESIS ON A SHOESTRING: IMPROVING THE EFFICIENCY OF LPCNET	8437
<i>Jean-Marc Valin, Umut Isik, Paris Smaragdis, Arvinth Krishnaswamy, Amazon, Canada</i>	
SPE-84.5: GENERALIZATION ABILITY OF MOS PREDICTION NETWORKS.....	8442
<i>Erica Cooper, Junichi Yamagishi, National Institute of Informatics, Japan; Wen-Chin Huang, Tomoki Toda, Nagoya University, Japan</i>	
SPE-84.6: ON THE INTERPLAY BETWEEN SPARSITY, NATURALNESS, INTELLIGIBILITY, AND PROSODY IN SPEECH SYNTHESIS	8447
<i>Cheng-I Lai, Yi-Lun Liao, Yung-Sung Chuang, Alexander Liu, James Glass, MIT CSAIL, United States of America; Erica Cooper, Junichi Yamagishi, National Institute of Informatics, Japan; Yang Zhang, Shiyu Chang, Kaizhi Qian, David Cox, MIT-IBM Watson AI Lab, United States of America</i>	
SPE-85: SIGN LANGUAGE AND LIP READING	
SPE-85.1: PHONOLOGY RECOGNITION IN AMERICAN SIGN LANGUAGE	8452
<i>Federico Tavella, Aphrodite Galata, Angelo Cangelosi, The University of Manchester, United Kingdom of Great Britain and Northern Ireland</i>	
SPE-85.2: SPATIO-TEMPORAL GRAPH CONVOLUTIONAL NETWORKS FOR CONTINUOUS SIGN LANGUAGE RECOGNITION	8457
<i>Maria Parelli, Petros Maragos, National Technical University of Athens, Greece; Katerina Papadimitriou, Gerasimos Potamianos, University of Thessaly, Greece; Georgios Pavlakos, University of California, United States of America</i>	
SPE-85.3: SENSORS TO SIGN LANGUAGE: A NATURAL APPROACH TO EQUITABLE COMMUNICATION	8462
<i>Thomas Fouts, University of Michigan, United States of America; Ali Hindy, Stanford University, United States of America; Chris Tanner, Harvard University, United States of America</i>	
SPE-85.4: ACCURATE AND RESOURCE-EFFICIENT LIPREADING WITH EFFICIENTNETV2 AND TRANSFORMERS	8467
<i>Alexandros Koumparoulis, Gerasimos Potamianos, University of Thessaly, Greece</i>	
SPE-85.5: TRAINING STRATEGIES FOR IMPROVED LIP-READING	8472
<i>Pingchuan Ma, Yujiang Wang, Stavros Petridis, Jie Shen, Maja Pantic, Imperial College London, United Kingdom of Great Britain and Northern Ireland</i>	
SPE-85.6: MULTISTREAM NEURAL ARCHITECTURES FOR CUED SPEECH RECOGNITION USING A PRE-TRAINED VISUAL FEATURE EXTRACTOR AND CONSTRAINED CTC DECODING	8477
<i>Sanjana Sankar, Denis Beutemps, Thomas Hueber, Centre National de la Recherche Scientifique, France</i>	
SPE-86: SPEECH PARALINGUISTICS: REPRESENTATION LEARNING	
SPE-86.1: MODELING OF PRE-TRAINED NEURAL NETWORK EMBEDDINGS	8482
LEARNED FROM RAW WAVEFORM FOR COVID-19 INFECTION DETECTION	
<i>Zohreh Mostaani, RaviShankar Prasad, Bogdan Vlasenko, Mathew Magimai Doss, Idiap Research Institute, Switzerland</i>	

SPE-86.2: DUAL ATTENTION POOLING NETWORK FOR RECORDING DEVICE CLASSIFICATION USING NEUTRAL AND WHISPERED SPEECH	8487
<i>Abinay Reddy Naini, Prasanta Kumar Ghosh, Indian Institute of Science, Bangalore, India; Bhavuk Singhal, Information Technology, Bundelkhand Institute of Engineering and Technology, Jhansi, India, India</i>	
SPE-86.3: ENTRAINMENT ANALYSIS FOR ASSESSMENT OF AUTISTIC SPEECH PROSODY USING BOTTLENECK FEATURES OF DEEP NEURAL NETWORK	8492
<i>Keiko Ochi, Kyoto University, Japan; Nobutaka Ono, Tokyo Metropolitan University, Japan; Keiho Owada, Miho Kuroda, Shigeki Sagayama, University of Tokyo, Japan; Hidenori Yamasue, Hamamatsu University School of Medicine, Japan</i>	
SPE-86.4: CUSTOMER SATISFACTION ESTIMATION USING UNSUPERVISED REPRESENTATION LEARNING WITH MULTI-FORMAT PREDICTION LOSS	8497
<i>Atsushi Ando, Yumiko Murata, Ryo Masumura, Satoshi Suzuki, Naoki Makishima, Takafumi Moriya, Takanori Ashihara, Hiroshi Sato, NTT Corporation, Japan</i>	
SPE-86.5: AUTOMATIC ASSESSMENT OF THE DEGREE OF CLINICAL DEPRESSION FROM SPEECH USING X-VECTORS	8502
<i>José Vicente Egas-López, University of Szeged, Hungary; Gábor Kiss, Sztahó David, Budapest University of Technology and Economics, Hungary; Gábor Gosztolya, MTA-SZTE Research Group on Artificial Intelligence, Hungary</i>	
SPE-86.6: AUTOMATIC DEPRESSION LEVEL ASSESSMENT FROM SPEECH BY LONG-TERM GLOBAL INFORMATION EMBEDDING	8507
<i>Ya Li, Beijing University of Posts and Telecommunications, China; Mingyue Niu, Ziping Zhao, Tianjin Normal University, China; Jianhua Tao, Institute of Automation, Chinese Academy of Sciences, China</i>	
SPE-87: SPEECH RECOGNITION: KNOWLEDGE TRANSFER AND CONTEXTUALIZATION	
SPE-87.1: KNOWLEDGE TRANSFER FROM LARGE-SCALE PRETRAINED LANGUAGE MODELS TO END-TO-END SPEECH RECOGNIZERS	8512
<i>Yotaro Kubo, Shigeki Karita, Michiel Bacchiani, Google, Japan</i>	
SPE-87.2: IMPROVING CTC-BASED SPEECH RECOGNITION VIA KNOWLEDGE TRANSFERRING FROM PRE-TRAINED LANGUAGE MODELS	8517
<i>Keqi Deng, Gaofeng Cheng, Ji Xu, Pengyuan Zhang, Institute of Acoustics, Chinese Academy of Sciences, China; Songjun Cao, Yike Zhang, Long Ma, Tencent Cloud Xiaowei, China</i>	
SPE-87.3: IMPROVING NON-AUTOREGRESSIVE END-TO-END SPEECH RECOGNITION WITH PRE-TRAINED ACOUSTIC AND LANGUAGE MODELS	8522
<i>Keqi Deng, Zehui Yang, Gaofeng Cheng, Pengyuan Zhang, Institute of Acoustics, Chinese Academy of Sciences, China; Shinji Watanabe, Carnegie Mellon University, United States of America; Yosuke Higuchi, Waseda University, Japan</i>	
SPE-87.4: KNOWLEDGE DISTILLATION FOR NEURAL TRANSDUCERS FROM LARGE SELF-SUPERVISED PRE-TRAINED MODELS	8527
<i>Xiaoyu Yang, Qiuji Li, Phil Woodland, University of Cambridge, United Kingdom of Great Britain and Northern Ireland</i>	
SPE-87.5: IMPROVING END-TO-END CONTEXTUAL SPEECH RECOGNITION WITH FINE-GRAINED CONTEXTUAL KNOWLEDGE SELECTION	8532
<i>Minglun Han, Shiyu Zhou, Bo Xu, Institute of Automation, Chinese Academy of Sciences, China; Linhao Dong, Zhenlin Liang, Meng Cai, Zejun Ma, Bytedance AI Lab, China</i>	
SPE-87.6: CONTEXTUAL ADAPTERS FOR PERSONALIZED SPEECH RECOGNITION IN NEURAL TRANSDUCERS	8537
<i>Kanthashree Mysore Sathyendra, Thejaswi Muniyappa, Feng-Ju Chang, Jing Liu, Jinru Su, Grant P. Strimel, Athanasios Mouchtaris, Siegfried Kunzmann, Amazon, United States of America</i>	

SPE-88: SPEECH ANALYSIS: PARALINGUISTICS

SPE-88.1: EMOTIONFLOW: CAPTURE THE DIALOGUE LEVEL EMOTION 8542 TRANSITIONS

Xiaohui Song, 1. Institute of Information Engineering, Chinese Academy of Sciences; 2. School of Cyber Security, University of Chinese Academy of Sciences; 3. Alibaba Group, China; Liangjun Zang, Songlin Hu, 1. Institute of Information Engineering, Chinese Academy of Sciences; 2. School of Cyber Security, University of Chinese Academy of Sciences, China; Rong Zhang, Longtao Huang, Alibaba Group, China

SPE-88.2: MULTIMODAL SENTIMENT ANALYSIS ON UNALIGNED SEQUENCES VIA 8547 HOLOGRAPHIC EMBEDDING

Yukun Ma, Bin Ma, Alibaba Group, Singapore

SPE-88.3: DISTRIBUTION LEARNING FOR AGE ESTIMATION FROM SPEECH..... 8552

Amruta Saraf, Elie Khoury, Pindrop, United States of America

SPE-88.4: DISPEECH: A SYNTHETIC TOY DATASET FOR SPEECH DISENTANGLING 8557

Olivier Zhang, Nicolas Gengembre, Olivier Le Blouch, Orange, France; Damien Lolive, IRISA, France

SPE-88.5: END-TO-END ASR-ENHANCED NEURAL NETWORK FOR ALZHEIMER'S 8562 DISEASE DIAGNOSIS

Jiancheng Gui, Yikai Li, Kai Chen, Joanna Siebert, Qingcai Chen, Harbin Institute of Technology (Shenzhen), China, China

SPE-88.6: A NOVEL SEQUENTIAL MONTE CARLO FRAMEWORK FOR PREDICTING 8567 AMBIGUOUS EMOTION STATES

Jingyao Wu, Vidhyasaharan Sethu, Eliathamby Ambikairajah, University of New South Wales, Australia; Ting Dang, University of Cambridge, United Kingdom of Great Britain and Northern Ireland

SPE-89: AUGMENTATION FOR SPEECH RECOGNITION

SPE-89.1: PHONE-INFORMED REFINEMENT OF SYNTHESIZED MEL 8572 SPECTROGRAM FOR DATA AUGMENTATION IN SPEECH RECOGNITION

Sei Ueno, Tatsuya Kawahara, Graduate School of Informatics, Kyoto University, Japan

SPE-89.2: LPC AUGMENT: AN LPC-BASED ASR DATA AUGMENTATION ALGORITHM 8577 FOR LOW AND ZERO-RESOURCE CHILDREN'S DIALECTS

Alexander Johnson, Ruchao Fan, Abeer Alwan, University of California - Los Angeles, United States of America; Robin Morris, Georgia State University, United States of America

SPE-89.3: TOWARDS BETTER META-INITIALIZATION WITH TASK AUGMENTATION 8582 FOR KINDERGARTEN-AGED SPEECH RECOGNITION

Yunzheng Zhu, Ruchao Fan, Abeer Alwan, University of California, Los Angeles, United States of America

SPE-89.4: UNSUPERVISED DATA SELECTION FOR SPEECH RECOGNITION WITH 8587 CONTRASTIVE LOSS RATIOS

Chanh Park, Rehan Ahmad, Thomas Hain, The University of Sheffield, United Kingdom of Great Britain and Northern Ireland

SPE-89.5: IMPORTANTAUG: A DATA AUGMENTATION AGENT FOR SPEECH..... 8592

Viet Anh Trinh, Hassan Kavaki, The City University of New York, United States of America; Michael Mandel, Brooklyn College, United States of America

SPE-89.6: INJECTING TEXT AND CROSS-LINGUAL SUPERVISION IN FEW-SHOT 8597 LEARNING FROM SELF-SUPERVISED MODELS

Matthew Wiesner, Human Language Technology Center of Excellence, United States of America; Desh Raj, Sanjeev Khudanpur, Johns Hopkins University, United States of America

SS-1: QUANTUM MACHINE LEARNING FOR SPEECH AND LANGUAGE PROCESSING

SS-1.1: WHEN BERT MEETS QUANTUM TEMPORAL CONVOLUTION LEARNING 8602 FOR TEXT CLASSIFICATION IN HETEROGENEOUS COMPUTING

Chao-Han Huck Yang, Jun Qi, Georgia Institute of Technology, United States of America; Samuel Yen-Chi Chen, Brookhaven National Laboratory, United States of America; Yu Tsao, Academia Sinica, Taiwan, Province of China; Pin-Yu Chen, IBM Research, United States of America

SS-1.2: MATCHING POINT SETS WITH QUANTUM CIRCUIT LEARNING..... 8607

Mohammadreza Noormandipour, Hanchen Wang, University of Cambridge, United Kingdom of Great Britain and Northern Ireland

SS-1.3: THE DAWN OF QUANTUM NATURAL LANGUAGE PROCESSING 8612

Riccardo Di Sipio, Ceridian HCM Inc., Italy; Jia-Hong Huang, Marcel Worring, University of Amsterdam, Netherlands; Samuel Yen-Chi Chen, Brookhaven National Laboratory, United States of America; Stefano Mangini, University of Pavia, Italy

SS-1.4: QUANTUM FEDERATED LEARNING WITH QUANTUM DATA 8617

Mahdi Chehimi, Walid Saad, Virginia Tech, United States of America

SS-1.5: QUANTUM LONG SHORT-TERM MEMORY 8622

Samuel Yen-Chi Chen, Shinjae Yoo, Yao-Lung L. Fang, Brookhaven National Laboratory, United States of America

SS-1.6: CLASSICAL-TO-QUANTUM TRANSFER LEARNING FOR SPOKEN COMMAND 8627 RECOGNITION BASED ON QUANTUM NEURAL NETWORKS

Jun Qi, Georgia Institute of Technology, United States of America; Javier Tejedor, Universidad San Pablo-CEU, CEU Universities, Spain

SS-2: SIGNAL PROCESSING AND MACHINE LEARNING FOR ENERGY HARVESTING COMMUNICATIONS

SS-2.1: WAVEFORM OPTIMIZATION FOR WIRELESS POWER TRANSFER WITH 8632 POWER AMPLIFIER AND ENERGY HARVESTER NON-LINEARITIES

Yumeng Zhang, Bruno Clerckx, Imperial College London, United Kingdom of Great Britain and Northern Ireland

SS-2.2: ECONOMICS OF SEMANTIC COMMUNICATION SYSTEM IN WIRELESS 8637 POWERED INTERNET OF THINGS

Zi Qin Liew, Yanyu Cheng, Wei Yang Bryan Lim, Alibaba-NTU Singapore Joint Research Institute, Nanyang Technological University, Singapore, Singapore; Dusit Niyato, School of Computer Science and Engineering, Nanyang Technological University, Singapore, Singapore; Chunyan Miao, Alibaba-NTU Singapore Joint Research Institute, Nanyang Technological University, Singapore and School of Computer Science and Engineering, Nanyang Technological University, Singapore, Singapore; Sumei Sun, Institute of Infocomm Research, Agency for Science, Technology and Research, Singapore, Singapore

SS-2.3: OPTIMAL RESOURCE ALLOCATION AND BEAMFORMING FOR TWO-USER 8642 MISO WPCNS FOR A NON-LINEAR CIRCUIT-BASED EH MODEL

Nikita Shanin, Moritz Garkisch, Robert Schober, Laura Cottatellucci, Friedrich-Alexander-Universität Erlangen-Nürnberg, Germany; Amelie Hagelauer, Fraunhofer EMFT Einrichtung für Mikrosysteme und Festkörper-Technologien; Technische Universität München, Germany

SS-2.4: PERFORMANCE OPTIMIZATION FOR WIRELESS SEMANTIC 8647 COMMUNICATIONS OVER ENERGY HARVESTING NETWORKS

Mingzhe Chen, H. Vincent Poor, Princeton University, United States of America; Yining Wang, Beijing University of Posts and Telecommunications, China

SS-2.5: DEEP LEARNING BASED PASSIVE BEAMFORMING FOR IRS-ASSISTED 8652 MONOSTATIC BACKSCATTER SYSTEMS

Sahar Idrees, Xiaolun Jia, Saud Khan, Salman Durrani, Xiangyun Zhou, Australian National University, Australia

SS-2.6: ON FEDERATED LEARNING WITH ENERGY HARVESTING CLIENTS	8657
<i>Cong Shen, University of Virginia, United States of America; Jing Yang, The Pennsylvania State University, United States of America; Jie Xu, University of Miami, United States of America</i>	

SS-3: MODEL-BASED DEEP LEARNING IN SIGNAL PROCESSING

SS-3.1: STRUCTURAL PRIOR MODELS FOR 3-D DEEP VESSEL SEGMENTATION.....	8662
<i>Xuelu Li, Vishal Monga, Pennsylvania State University, University Park, United States of America; Raja Bala, Amazon Research, United States of America</i>	

SS-3.2: EXPECTATION CONSISTENT PLUG-AND-PLAY FOR MRI	8667
<i>Saurav Shastri, Rizwan Ahmad, Philip Schniter, Ohio State, United States of America; Christopher Metzler, Univ. Maryland, United States of America</i>	

SS-3.3: INVERSE IMAGING WITH GENERATIVE PRIORS VIA LANGEVIN DYNAMICS.....	8672
<i>Thanh Nguyen, Unaffiliated, United States of America; Gauri Jagatap, Chinmay Hegde, New York University, United States of America</i>	

SS-3.4: CNN-AIDED FACTOR GRAPHS WITH ESTIMATED MUTUAL INFORMATION	8677
FEATURES FOR SEIZURE DETECTION	
<i>Bahareh Salafian, Nariman Farsad, Ryerson University, Canada; Eyal Fishel Ben-Knaan, Nir Shlezinger, Ben-Gurion University of the Negev, Israel; Sandrine de Ribaupierre, University of Western Ontario, Canada</i>	

SS-3.5: UNFOLDING MODEL-BASED BEAMFORMING FOR HIGH QUALITY	8682
ULTRASOUND IMAGING	
<i>Christopher Khan, Brett Byram, Vanderbilt University, United States of America; Ruud van Sloun, Eindhoven University of Technology, Netherlands</i>	

SS-3.6: OPTIMIZATION GUARANTEES FOR ISTA AND ADMM BASED UNFOLDED	8687
NETWORKS	
<i>Wei Pu, Miguel Rodrigues, UCL, United Kingdom of Great Britain and Northern Ireland; Yonina Eldar, Weizmann institute of science, Israel</i>	

SS-4: HUMAN CENTRIC ACTIVE NOISE CONTROL APPLICATIONS

SS-4.1: INTEGRATION OF ANOMALY MACHINE SOUND DETECTION INTO ACTIVE	8692
NOISE CONTROL TO SHAPE THE RESIDUAL SOUND	
<i>Chuang Shi, Mengjie Huang, Huitian Jiang, Huiyong Li, University of Electronic Science and Technology of China, China</i>	

SS-4.2: DUAL ACTIVE NOISE CONTROL WITH COMMON SENSORS.....	8697
<i>Ryosuke Okajima, Yoshinobu Kajikawa, Kansai University, Japan; Kohei Oto, Nitto Denko Corporation, Japan</i>	

SS-4.3: A HYBRID APPROACH TO COMBINE WIRELESS AND EARCUP	8702
MICROPHONES FOR ANC HEADPHONES WITH ERROR SEPARATION MODULE	
<i>Xiaoyi Shen, Dongyuan Shi, Woon-Seng Gan, Nanyang Technological University, Singapore</i>	

SS-4.4: SPATIAL ACTIVE NOISE CONTROL WITH THE REMOTE MICROPHONE	8707
TECHNIQUE: AN APPROACH WITH A MOVING HIGHER ORDER MICROPHONE	
<i>Huiyuan Sun, Jihui Zhang, Thushara Abhayapala, Prasanga Samarasinghe, The Australian National University, Australia</i>	

SS-4.5: ROBUST PRESSURE MATCHING WITH ATF PERTURBATION CONSTRAINTS	8712
FOR SOUND FIELD CONTROL	
<i>Junqing Zhang, Wen Zhang, Lijun Zhang, Jingdong Chen, Center of Intelligent Acoustics and Immersive Communications, School of Marine Science and Technology, Northwestern Polytechnical University, China; Liming Shi, Mads Græsbøll Christensen, Audio Analysis Lab, CREATE, Aalborg University, China</i>	

SS-4.6: OPTIMIZATION OF A FIXED VIRTUAL SENSING FEEDBACK AND CONTROLLER FOR IN-EAR HEADPHONES WITH MULTIPLE LOUDSPEAKERS	8717
<i>Piero Rivera Benois, Simon Doclo, Signal Processing Group, University of Oldenburg, Oldenburg, Germany, Germany; Reinhild Roden, Matthias Blau, Institut fuer Hoertechnik und Audiologie, Jade Hochschule, Oldenburg, Germany, Germany</i>	
SS-5: INTEGRATED SENSING AND COMMUNICATIONS FOR FUTURE CELLULAR AND VEHICULAR NETWORKS	
SS-5.1: ON THE POTENTIAL OF SPATIALLY-SPREAD ORTHOGONAL TIME FREQUENCY SPACE MODULATION FOR ISAC TRANSMISSIONS	8722
<i>Shuangyang Li, Jinhong Yuan, University of New South Wales, Australia; Weijie Yuan, Southern University of Science and Technology, China; Giuseppe Caire, Technische Universität Berlin, China</i>	
SS-5.2: SENSING-ASSISTED BEAM TRACKING IN V2I NETWORKS: EXTENDED TARGET CASE	8727
<i>Zhen Du, Zenghui Zhang, Shanghai Jiao Tong University, China; Fan Liu, Southern University of Science and Technology, China</i>	
SS-5.3: TRANSMIT BEAMFORMING WITH FIXED COVARIANCE FOR INTEGRATED MIMO RADAR AND MULTIUSER COMMUNICATIONS	8732
<i>Xiang Liu, Tianyao Huang, Yimin Liu, Tsinghua University, China; Yonina Eldar, Weizmann institute of Science, Israel</i>	
SS-5.4: SAFEGUARDING UAV NETWORKS THROUGH INTEGRATED SENSING, JAMMING, AND COMMUNICATIONS	8737
<i>Zhiqiang Wei, Robert Schober, Friedrich-Alexander University Erlangen-Nuremberg, Germany; Fan Liu, Southern University of Science and Technology, China; Derrick Wing Kwan Ng, University of New South Wales, Australia</i>	
SS-5.5: EVALUATION OF ORTHOGONAL CHIRP DIVISION MULTIPLEXING FOR AUTOMOTIVE INTEGRATED SENSING AND COMMUNICATIONS	8742
<i>Sangeeta Bhattacharjee, Chandra Murthy, Indian Institute of Science, Bangalore, India; Kumar Vijay Mishra, United States CDC Army Research Laboratory, United States of America; Ramesh Annavajjala, University of Massachusetts, United States of America</i>	
SS-5.6: INTEGRATED SENSING AND COMMUNICATIONS VIA 5G NR WAVEFORM: PERFORMANCE ANALYSIS	8747
<i>Yuanhao Cui, Xiaojun Jing, Junsheng Mu, Beijing University of Posts and Telecommunications, China</i>	
SS-6: FRONTIERS OF FEDERATED LEARNING: APPLICATIONS, CHALLENGES, AND OPPORTUNITIES	
SS-6.1: FEDERATED LEARNING CHALLENGES AND OPPORTUNITIES: AN OUTLOOK	8752
<i>Jie Ding, Eric Tramel, Anit Kumar Sahu, Shuang Wu, Salman Avestimehr, Tao Zhang, Amazon, United States of America</i>	
SS-6.2: ENABLING ON-DEVICE TRAINING OF SPEECH RECOGNITION MODELS WITH FEDERATED DROPOUT	8757
<i>Dhruv Guliani, Lillian Zhou, Changwan Ryu, Tien-Ju Yang, Harry Zhang, Yonghui Xiao, Françoise Beaufays, Giovanni Motta, Google Inc, United States of America</i>	
SS-6.3: ADAPTIVE NODE PARTICIPATION FOR STRAGGLER-RESILIENT FEDERATED LEARNING	8762
<i>Amirhossein Reisizadeh, MIT, United States of America; Isidoros Tziotis, Aryan Mokhtari, UT Austin, United States of America; Hamed Hassani, UPenn, United States of America; Ramtin Pedarsani, UC Santa Barbara, United States of America</i>	

SS-6.4: LEARNINGS FROM FEDERATED LEARNING IN THE REAL WORLD	8767
<i>Christophe Dupuy, Tanya Roosta, Leo Long, Clement Chung, Rahul Gupta, Amazon Alexa AI, United States of America; Salman Avestimehr, Amazon Alexa AI ; University of Southern California (USC), United States of America</i>	
SS-6.5: A DYNAMIC REWEIGHTING STRATEGY FOR FAIR FEDERATED LEARNING.....	8772
<i>Zhiyuan Zhao, Gauri Joshi, Carnegie Mellon University, United States of America</i>	
SS-6.6: OVER-THE-AIR PERSONALIZED FEDERATED LEARNING.....	8777
<i>Hasin Us Sami, Basak Guler, University of California, Riverside, United States of America</i>	
 SS-7: LEARNING-BASED ACOUSTIC SIGNAL PROCESSING FOR STAY-AT-HOME APPLICATIONS IN PANDEMIC TIMES	
SS-7.1: DNN BASED MULTIFRAME SINGLE-CHANNEL NOISE REDUCTION	8782
FILTERS	
<i>Ningning Pan, Jingdong Chen, Northwestern Polytechnical University, China; Jacob Benesty, University of Quebec, Canada</i>	
SS-7.2: LEARNING-BASED PERSONAL SPEECH ENHANCEMENT FOR	8787
TELECONFERENCING BY EXPLOITING SPATIAL-SPECTRAL FEATURES	
<i>Yicheng Hsu, Yonghan Lee, Mingsian Bai, National Tsing Hua University, Taiwan</i>	
SS-7.3: MANIFOLD LEARNING-SUPPORTED ESTIMATION OF RELATIVE TRANSFER	8792
FUNCTIONS FOR SPATIAL FILTERING	
<i>Andreas Brendel, Johannes Zeitler, Walter Kellermann, Friedrich-Alexander University Erlangen-Nuremberg, Germany</i>	
SS-7.4: AUDIO SIGNAL PROCESSING FOR TELEPRESENCE BASED ON WEARABLE	8797
ARRAY IN NOISY AND DYNAMIC SCENES	
<i>Hanan Beit-On, Moti Lugasi, Lior Madmoni, Boaz Rafaely, Ben-Gurion University, Israel; Anjali Menon, Anurag Kumar, Jacob Donley, Vladimir Tourbabin, Facebook, United States of America</i>	
SS-7.5: A MULTI-TASK LEARNING METHOD FOR WEAKLY SUPERVISED SOUND	8802
EVENT DETECTION	
<i>Sichen Liu, Feiran Yang, Fang Kang, Jun Yang, Institute of Acoustics, Chinese Academy of Sciences; University of Chinese Academy of Sciences, China</i>	
SS-7.6: LOW RESOURCES ONLINE SINGLE-MICROPHONE SPEECH	8807
ENHANCEMENT WITH HARMONIC EMPHASIS	
<i>Nir Raviv, Ofer Schwartz, CEVA-DSP, Israel; Sharon Gannot, Bar-Ilan University, Israel</i>	
 SS-8: MACHINE LEARNING IN BEYOND 5G WIRELESS NETWORKS	
SS-8.1: DEEP LEARNING FOR LOCATION BASED BEAMFORMING WITH NLOS	8812
CHANNELS	
<i>Luc Le Magoarou, Stéphane Paquelet, bcom, France; Taha Yassine, Matthieu Crussière, bcom and INSA Rennes, France</i>	
SS-8.2: PREDICTING FLAT-FADING CHANNELS VIA META-LEARNED CLOSED-FORM	8817
LINEAR FILTERS AND EQUILIBRIUM PROPAGATION	
<i>Sangwoo Park, Osvaldo Simeone, King's College London, United Kingdom of Great Britain and Northern Ireland</i>	
SS-8.3: DEEP-LEARNING-ASSISTED CONFIGURATION OF RECONFIGURABLE	8822
INTELLIGENT SURFACES IN DYNAMIC RICH-SCATTERING ENVIRONMENTS	
<i>Kyriakos Stylianopoulos, George Alexandropoulos, National and Kapodistrian University of Athens, Greece; Nir Shlezinger, Ben-Gurion University of the Negev, Israel; Philipp del Hougne, Univ Rennes, France</i>	

SS-8.4: SUPERVISED LEARNING BASED SPARSE CHANNEL ESTIMATION FOR RIS AIDED COMMUNICATIONS	8827
<i>Dilin Dampahalage, K. B. Shashika Manosha, Nandana Rajatheva, Matti Latva-aho, University of Oulu, Finland</i>	
SS-8.5: GOAL-ORIENTED COMMUNICATION FOR EDGE LEARNING BASED ON THE INFORMATION BOTTLENECK	8832
<i>Francesco Pezone, Sergio Barbarossa, Paolo Di Lorenzo, Sapienza University of Rome, Italy</i>	
SS-9: PROCESSING HIGHER ORDER NETWORK DATA	
SS-9.1: HYPERGRAPHS WITH EDGE-DEPENDENT VERTEX WEIGHTS: SPECTRAL CLUSTERING BASED ON THE 1-LAPLACIAN	8837
<i>Yu Zhu, Boning Li, Santiago Segarra, Rice University, United States of America</i>	
SS-9.2: CAUSAL LINEAR TOPOLOGICAL FILTERS OVER A 2-SIMPLEX	8842
<i>Georg Essl, University of Wisconsin - Milwaukee, United States of America</i>	
SS-9.3: SIMPLICIAL CONVOLUTIONAL NEURAL NETWORKS	8847
<i>Maosheng Yang, Elvin Isufi, Geert Leus, Delft University of Technology, Netherlands</i>	
SS-9.4: SIGNAL PROCESSING ON CELL COMPLEXES	8852
<i>T. Mitchell Roddenberry, Rice University, United States of America; Michael T. Schaub, RWTH Aachen University, Germany; Mustafa Hajij, Santa Clara University, United States of America</i>	
SS-9.5: ROBUST SIGNAL PROCESSING OVER SIMPLICIAL COMPLEXES	8857
<i>Stefania Sardellitti, Sergio Barbarossa, Sapienza University of Rome, Italy</i>	
SS-10: SIGNAL PROCESSING AND NEURAL APPROACHES FOR SOUNDSCAPES (SINAPS)	
SS-10.1: CONFORMER-BASED SELF-SUPERVISED LEARNING FOR NON-SPEECH AUDIO TASKS	8862
<i>Sangeeta Srivastava, The Ohio State University, United States of America; Yun Wang, Andros Tjandra, Anurag Kumar, Chunxi Liu, Kritika Singh, Yatharth Saraf, Meta, United States of America</i>	
SS-10.2: UNSUPERVISED AUDIO-CAPTION ALIGNING LEARNS CORRESPONDENCES BETWEEN INDIVIDUAL SOUND EVENTS AND TEXTUAL PHRASES	8867
<i>Huang Xie, Okko Räsänen, Konstantinos Drossos, Tuomas Virtanen, Tampere University, Finland</i>	
SS-10.3: SPATIAL DATA AUGMENTATION WITH SIMULATED ROOM IMPULSE RESPONSES FOR SOUND EVENT LOCALIZATION AND DETECTION	8872
<i>Yuichiro Koyama, Masafumi Takahashi, Kazuki Shimada, Naoya Takahashi, Emiru Tsunoo, Shusuke Takahashi, Yuki Mitsufuji, Sony Group Corporation, Japan; Kazuhide Shigemi, The University of Tokyo, Japan</i>	
SS-10.4: POLYPHONIC AUDIO EVENT DETECTION: MULTI-LABEL OR MULTI-CLASS MULTI-TASK CLASSIFICATION PROBLEM?	8877
<i>Huy Phan, Queen Mary University of London, United Kingdom of Great Britain and Northern Ireland; Thi Ngoc Tho Nguyen, Nanyang Technological University, Singapore; Philipp Koch, Alfred Mertins, University of Lübeck, Germany</i>	
SS-10.5: DIVERSE AUDIO CAPTIONING VIA ADVERSARIAL TRAINING	8882
<i>Xinhao Mei, Xubo Liu, Jianyuan Sun, Mark Plumbley, Wenwu Wang, University of Surrey, United Kingdom of Great Britain and Northern Ireland</i>	
SS-10.6: PROBABLY PLEASANT? A NEURAL-PROBABILISTIC APPROACH TO AUTOMATIC MASKER SELECTION FOR URBAN SOUNDSCAPE AUGMENTATION	8887
<i>Kenneth Ooi, Karn N. Watcharasupat, Bhan Lam, Zhen-Ting Ong, Woon-Seng Gan, Nanyang Technological University, Singapore</i>	

SS-11: GRAPH-BASED DEEP LEARNING FOR COMMUNICATIONS

SS-11.1: USER SCHEDULING USING GRAPH NEURAL NETWORKS FOR 8892 **RECONFIGURABLE INTELLIGENT SURFACE ASSISTED MULTIUSER DOWNLINK** **COMMUNICATIONS**

Zhongze Zhang, Tao Jiang, Wei Yu, University of Toronto (St. George), Canada

SS-11.2: SYMBOL-LEVEL ONLINE CHANNEL TRACKING FOR DEEP RECEIVERS..... 8897 *Ron Aharon Finish, Yoav Cohen, Tomer Raviv, Nir Shlezinger, Ben-Gurion University of the Negev, Israel*

SS-11.3: DELAY-ORIENTED DISTRIBUTED SCHEDULING USING GRAPH NEURAL 8902 **NETWORKS**

Zhongyuan Zhao, Santiago Segarra, Rice University, United States of America; Gunjan Verma, Ananthram Swami, US Army's DEVCOM Army Research Laboratory, United States of America

SS-11.4: FLOWDT: A FLOW-AWARE DIGITAL TWIN FOR COMPUTER NETWORKS 8907 *Miquel Ferriol-Galmés, Pere Barlet-Ros, Albert Cabellos-Aparicio, Universitat Politècnica de Catalunya, Spain;* *Xiangle Cheng, University of Exeter, United Kingdom of Great Britain and Northern Ireland; Xiang Shi, Shihan Xiao,* *Huawei Technologies, China*

SS-11.5: STABLE AND TRANSFERABLE WIRELESS RESOURCE ALLOCATION 8912 **POLICIES VIA MANIFOLD NEURAL NETWORKS**

Zhiyang Wang, Luana Ruiz, Alejandro Ribeiro, University of Pennsylvania, United States of America; Mark Eisen,
Intel Labs, United States of America

SS-12: TOWARDS NEUROMORPHIC MACHINE INTELLIGENCE: SPIKE-BASED **REPRESENTATION, LEARNING, AND SIGNAL PROCESSING**

SS-12.1: MOTIF-TOPOLOGY AND REWARD-LEARNING IMPROVED SPIKING 8917 **NEURAL NETWORK FOR EFFICIENT MULTI-SENSORY INTEGRATION**

Shuncheng Jia, Tielin Zhang, Hongxing Liu, Bo Xu, Institute of Automation, Chinese Academy of Sciences, China;
Ruichen Zuo, School of Information and Electronics, Beijing Institute of Technology, China

SS-12.2: EVENT-BASED MULTIMODAL SPIKING NEURAL NETWORK WITH 8922 **ATTENTION MECHANISM**

Qianhui Liu, Dong Xing, Lang Feng, Huajin Tang, Gang Pan, Zhejiang University, China

SS-12.3: GRADUAL SURROGATE GRADIENT LEARNING IN DEEP SPIKING NEURAL 8927 **NETWORKS**

Yi Chen, Silin Zhang, Shiyu Ren, Hong Qu, University of Electronic Science and Technology of China, China

SS-12.4: AXONAL DELAY AS A SHORT-TERM MEMORY FOR FEED FORWARD DEEP 8932 **SPIKING NEURAL NETWORKS**

*Pengfei Sun, Dick Botteldooren, Ghent University, Belgium; Longwei Zhu, Institute for Infocomm Research, A*STAR,*
Singapore

SS-12.5: LOW PRECISION LOCAL LEARNING FOR HARDWARE-FRIENDLY 8937 **NEUROMORPHIC VISUAL RECOGNITION**

*Jyotibdhya Acharya, Laxmi R Iyer, Wenyu Jiang, Institute of Infocomm Research, A*STAR, Singapore*

SS-12.6: A HYBRID LEARNING FRAMEWORK FOR DEEP SPIKING NEURAL 8942 **NETWORKS WITH ONE-SPIKE TEMPORAL CODING**

Jiadong Wang, Jibin Wu, Malu Zhang, Qi Liu, Haizhou Li, National University of Singapore, Singapore

SS-13: SEMANTIC AND INTERPRETABLE ANALYSIS IN MULTIMEDIA FORENSICS

SS-13.1: A-PIXELHOP: A GREEN, ROBUST AND EXPLAINABLE FAKE-IMAGE DETECTOR 8947

Yao Zhu, Xinyu Wang, Hong-Shuo Chen, C.-C. Jay Kuo, University of Southern California, United States of America; Ronald Salloum, California State University, San Bernardino, United States of America

SS-13.2: EXPLAINABLE FACT-CHECKING THROUGH QUESTION ANSWERING 8952

Jing Yang, Didier Vega-Oliveros, Anderson Rocha, Recod.ai, University of Campinas, Brazil; Taís Seibt, Universidade do Vale do Rio dos Sinos, Brazil

SS-13.3: DEEP VIDEO INPAINTING LOCALIZATION USING SPATIAL AND TEMPORAL TRACES 8957

Shujin Wei, Haodong Li, Jiwu Huang, Shenzhen University, China

SS-13.4: DEEPFAKE SPEECH DETECTION THROUGH EMOTION RECOGNITION: A SEMANTIC APPROACH 8962

Emanuele Conti, Davide Salvi, Clara Borrelli, Paolo Bestagini, Fabio Antonacci, Augusto Sarti, Stefano Tubaro, Politecnico di Milano, Italy; Brian Hosler, Matthew C. Stamm, Drexel University, United States of America

SS-13.5: TEXT-IMAGE DE-CONTEXTUALIZATION DETECTION USING VISION-LANGUAGE MODELS 8967

Mingzhen Huang, Shan Jia, Siwei Lyu, University at Buffalo, United States of America; Ming-Ching Chang, University at Albany, United States of America

SS-13.6: CUSTOM ATTRIBUTION LOSS FOR IMPROVING GENERALIZATION AND INTERPRETABILITY OF DEEPFAKE DETECTION 8972

Pavel Korshunov, Anubhav Jain, Sebastien Marcel, Idiap Research Institute, Switzerland

SS-14: GROUP THEORY FOR SENSING, IMAGING, AND LEARNING APPLICATIONS

SS-14.1: SPARSE MULTI-REFERENCE ALIGNMENT: SAMPLE COMPLEXITY AND COMPUTATIONAL HARDNESS 8977

Tamir Bendory, Tel Aviv University, Israel; Oscar Mickelin, Amit Singer, Princeton University, United States of America

SS-14.2: GRASSMANNIAN DIMENSIONALITY REDUCTION USING TRIPLET MARGIN LOSS FOR UME CLASSIFICATION OF 3D POINT CLOUDS 8982

Yuval Haitman, Joseph Francos, Ben-Gurion University, Israel; Louis Scharf, Colorado State University, United States of America

SS-14.3: A NOTE ON TOTALLY SYMMETRIC EQUI-ISOCLINIC TIGHT FUSION FRAMES 8987

Matthew Fickus, Air Force Institute of Technology, United States of America; Joseph Iverson, Iowa State University, United States of America; John Jasper, South Dakota State University, United States of America; Dustin Mixon, The Ohio State University, United States of America

SS-14.4: A SIMPLE FORMULA FOR THE MOMENTS OF UNITARILY INVARIANT MATRIX DISTRIBUTIONS 8992

Stephen D. Howard, Defence Science and Technology Group, Australia; Ali Pezeshki, Colorado State University, United States of America

SS-15: COMPUTER AUDITION FOR HEALTHCARE (CA4H)

SS-15.1: FUSION OF MODULATION SPECTRAL AND SPECTRAL FEATURES WITH SYMPTOM METADATA FOR IMPROVED SPEECH-BASED COVID-19 DETECTION 8997

Yi Zhu, Tiago Falk, Institut national de la recherche scientifique, Canada

SS-15.2: AN OVERVIEW OF THE FIRST ICASSP SPECIAL SESSION ON COMPUTER AUDITION FOR HEALTHCARE	9002
<i>Kun Qian, Beijing Institute of Technology, China; Tanja Schultz, University of Bremen, Germany; Björn Schuller, Imperial College London, United Kingdom of Great Britain and Northern Ireland</i>	
SS-15.3: A GLANCE-AND-GAZE NETWORK FOR RESPIRATORY SOUND CLASSIFICATION	9007
<i>Shuai Yu, Yiwei Ding, Wei Li, Fudan University, China; Kun Qian, Bin Hu, Beijing Institute of Technology, China; Bjoern Schuller, Imperial College London, United Kingdom of Great Britain and Northern Ireland</i>	
SS-15.4: INTERNET STREAMING AUDIO BASED SPEECH RECEPTION THRESHOLD MEASUREMENT IN COCHLEAR IMPLANT USERS	9012
<i>Xi Chen, Kang Ouyang, Mingyue Shi, Huali Zhou, Nengheng Zheng, Shenzhen University, China; Yefei Mo, Qinglin Meng, South China University of Technology, China; Yupeng Shi, Wei Xiao, Shidong Shang, Tencent, China</i>	
SS-15.5: A DOMAIN TRANSFER BASED DATA AUGMENTATION METHOD FOR AUTOMATED RESPIRATORY CLASSIFICATION	9017
<i>Zijie Wang, Zhao Wang, Zhejiang University, China</i>	
 SS-16: SECURITY- AND PRIVACY-PRESERVING SIGNAL PROCESSING TECHNIQUES FOR EMERGING APPLICATIONS	
SS-16.1: PHYSICAL LAYER ANONYMOUS COMMUNICATIONS: AN ANONYMITY ENTROPY ORIENTED PRECODING DESIGN	9022
<i>Zhongxiang Wei, Tongji University, China; Christos Masouros, University College London, United Kingdom of Great Britain and Northern Ireland; Sumei Sun, Institute for Infocomm Research (I2R), Singapore</i>	
SS-16.2: FEDERATED STOCHASTIC GRADIENT DESCENT BEGETS SELF-INDUCED MOMENTUM	9027
<i>Howard Yang, Zuozhu Liu, Zhejiang University, China; Yaru Fu, The Open University of Hong Kong, China; Tony Quek, Singapore University of Technology and Design, Singapore; Vincent Poor, Princeton University, United States of America</i>	
SS-16.3: ADVERSARIAL LEARNING IN TRANSFORMER BASED NEURAL NETWORK IN RADIO SIGNAL CLASSIFICATION	9032
<i>Lu Zhang, Sangarapillai Lambotharan, Gan Zheng, Loughborough University, United Kingdom of Great Britain and Northern Ireland</i>	
SS-16.4: OPTM3SEC: OPTIMIZING MULTICAST IRS-AIDED MULTIAN TENNA DFRC SECRECY CHANNEL WITH MULTIPLE EAVESDROPPERS	9037
<i>Kumar Vijay Mishra, United States Army Research Laboratory, United States of America; Arpan Chattopadhyay, Siddharth Sankar Acharjee, Indian Institute of Technology, India; Athina Petropulu, Rutgers - The State University of New Jersey, United States of America</i>	
SS-16.5: PRIVACY-ENHANCING APPLIANCE FILTERING FOR SMART METERS	9042
<i>Ramana Reddy Avula, Tobias Oechtering, KTH Royal Institute of Technology, Sweden</i>	
SS-16.6: ADVERSARIAL LINEAR QUADRATIC REGULATOR UNDER FALSIFIED ACTIONS	9047
<i>Chenglong Sun, Zuxing Li, Chao Wang, Tongji University, China</i>	
 SS-17: PROVABLE TENSOR AND MATRIX METHODS FOR SENSING AND LEARNING	
SS-17.1: COMMUNICATION-EFFICIENT DISTRIBUTED MAX-VAR GENERALIZED CCA VIA ERROR FEEDBACK-ASSISTED QUANTIZATION	9052
<i>Sagar Shrestha, Xiao Fu, Oregon State University, United States of America</i>	

SS-17.2: PROVABLE SECOND-ORDER RIEMANNIAN GAUSS-NEWTON METHOD	9057
FOR LOW-RANK TENSOR ESTIMATION	
<i>Yuetian Luo, University of Wisconsin-Madison, United States of America; Qin Ma, Ohio State University, United States of America; Chi Zhang, Indiana University, United States of America; Anru Zhang, Duke University, United States of America</i>	
SS-17.3: BOUNDED SIMPLEX-STRUCTURED MATRIX FACTORIZATION	9062
<i>Olivier Vu Thanh, Nicolas Gillis, Fabian Lecron, Université de Mons, Belgium</i>	
SS-17.4: CPD COMPUTATION VIA RECURSIVE EIGENSPACE DECOMPOSITIONS	9067
<i>Eric Evert, Michiel Vandecappelle, Lieven De Lathauwer, KU Leuven, Belgium</i>	
SS-17.5: ACCELERATING ILL-CONDITIONED ROBUST LOW-RANK TENSOR REGRESSION	9072
<i>Tian Tong, Yuejie Chi, Carnegie Mellon University, United States of America; Cong Ma, University of Chicago, United States of America</i>	
SS-17.6: ADA-JSR: SAMPLE EFFICIENT ADAPTIVE JOINT SUPPORT RECOVERY FROM EXTREMELY COMPRESSED MEASUREMENT VECTORS	9077
<i>Sina Shahsavari, Pulak Sarangi, Mehmet Hucumenoglu, Piya Pal, University of California San Diego, United States of America</i>	
SS-18: SIGNAL PROCESSING AND MACHINE LEARNING FOR RESPIRATORY DISEASE ANALYSIS	
SS-18.1: END-TO-END NETWORK BASED ON TRANSFORMER FOR AUTOMATIC DETECTION OF COVID-19	9082
<i>Cong Cai, Bin Liu, Jianhua Tao, Zhengkun Tian, Kexin Wang, Institute of Automation, Chinese Academy of Sciences, China; Jiahao Lu, Tianjin Normal University, China</i>	
SS-18.2: PROTOTYPE LEARNING FOR INTERPRETABLE RESPIRATORY SOUND ANALYSIS	9087
<i>Zhao Ren, Thanh Tam Nguyen, Wolfgang Nejdl, Leibniz Universität Hannover, Germany</i>	
SS-18.3: CONVOLUTIONAL TRANSFORMER WITH ADAPTIVE POSITION EMBEDDING FOR COVID-19 DETECTION FROM COUGH SOUNDS	9092
<i>Tianhao Yan, Harbin Engineering University and University of Augsburg, China; Hao Meng, Harbin Engineering University, China; Shuo Liu, University of Augsburg, Germany; Emilia Parada-Cabaleiro, Johannes Kepler University Linz, Austria; Zhao Ren, Leibniz Universität Hannover, Germany; Björn W. Schuller, University of Augsburg and Imperial College London, Germany</i>	
SS-18.4: DETECTION OF COPD EXACERBATION FROM SPEECH: COMPARISON OF ACOUSTIC FEATURES AND DEEP LEARNING BASED SPEECH BREATHING MODELS	9097
<i>Venkata Srikanth Nallanthighal, Philips Research Eindhoven and Radboud University, Nijmegen, Netherlands; Aki Härmä, Philips Research Eindhoven, Netherlands; Helmer Strik, Radboud University, Nijmegen, Netherlands</i>	
SS-18.5: AUTOMATIC RESPIRATORY SOUND CLASSIFICATION VIA MULTI-BRANCH TEMPORAL CONVOLUTIONAL NETWORK	9102
<i>Ziping Zhao, Zhen Gong, Mingyue Niu, Jiali Ma, Tianjin Normal University, China; Haishuai Wang, Fairfield University, United States of America; Zixing Zhang, Imperial College London, United Kingdom of Great Britain and Northern Ireland; Ya Li, Beijing University of Posts and Telecommunications, China</i>	
CHAL-1: ACOUSTIC ECHO CANCELLATION	
CHAL-1.1: ICASSP 2022 ACOUSTIC ECHO CANCELLATION CHALLENGE	9107
<i>Ross Cutler, Ando Saabas, Tanel Parnamaa, Marju Purin, Hannes Gamper, Sebastian Braun, Karsten Sørensen, Robert Aichner, Microsoft, United States of America</i>	

CHAL-1.2: A DEEP HIERARCHICAL FUSION NETWORK FOR FULLBAND ACOUSTIC ECHO CANCELLATION	9112
<i>Haoran Zhao, Nan Li, Runqiang Han, Lianwu Chen, Xiguang Zheng, Chen Zhang, Liang Guo, Bing Yu, Kuaishou Technology, China</i>	
CHAL-1.3: EXPLORE RELATIVE AND CONTEXT INFORMATION WITH TRANSFORMER FOR JOINT ACOUSTIC ECHO CANCELLATION AND SPEECH ENHANCEMENT	9117
<i>Xingwei Sun, Chenbin Cao, Qinglong Li, Linzhang Wang, Fei Xiang, Xiaomi Corporation, China</i>	
CHAL-1.4: MULTI-SCALE TEMPORAL FREQUENCY CONVOLUTIONAL NETWORK WITH AXIAL ATTENTION FOR SPEECH ENHANCEMENT	9122
<i>Guochang Zhang, Libiao Yu, Chunliang Wang, Jianqiang Wei, Baidu, China</i>	
CHAL-1.5: MULTI-TASK DEEP RESIDUAL ECHO SUPPRESSION WITH ECHO-AWARE LOSS	9127
<i>Shimin Zhang, Jiayao Sun, Yihui Fu, Lei Xie, Northwestern Polytechnical University, China; Ziteng Wang, Biao Tian, Qiang Fu, Personal, China</i>	
CHAL-1.6: MULTI-SCALE REFINEMENT NETWORK BASED ACOUSTIC ECHO CANCELLATION	9132
<i>Fan Cui, Liyong Guo, Wenfeng Li, Peng Gao, Yujun Wang, Xiaomi Inc, China</i>	
CHAL-2: AUDIO-VISUAL OBJECT CLASSIFICATION FOR HUMAN-ROBOT COLLABORATION	
CHAL-2.1: AUDIO-VISUAL OBJECT CLASSIFICATION FOR HUMAN-ROBOT COLLABORATION	9137
<i>Alessio Xompero, Yik Lung Pang, Andrea Cavallaro, Queen Mary University of London, United Kingdom of Great Britain and Northern Ireland; Timothy Patten, University of Technology Sydney, Australia; Ahalya Prabhakar, École Polytechnique Fédérale de Lausanne, Switzerland; Berk Calli, Worcester Polytechnic Institute, United States of America</i>	
CHAL-2.2: SHARED TRANSFORMER ENCODER WITH MASK-BASED 3D MODEL ESTIMATION FOR CONTAINER MASS ESTIMATION	9142
<i>Tomoya Matsubara, Seitaro Otsuki, Yuiga Wada, Haruka Matsuo, Takumi Komatsu, Yui Iioka, Komei Sugiura, Hideo Saito, Keio University, Japan</i>	
CHAL-2.3: IMPROVING GENERALIZATION OF DEEP NETWORKS FOR ESTIMATING PHYSICAL PROPERTIES OF CONTAINERS AND FILLINGS	9147
<i>Hengyi Wang, University College London, United Kingdom of Great Britain and Northern Ireland; Chaoran Zhu, Changjae Oh, Queen Mary University of London, China; Ziyin Ma, University of Birmingham, United Kingdom of Great Britain and Northern Ireland</i>	
CHAL-2.4: CONTAINER LOCALISATION AND MASS ESTIMATION WITH AN RGB-D CAMERA	9152
<i>Tommaso Apicella, Giulia Slavic, Edoardo Ragusa, Paolo Gastaldo, Lucio Marcenaro, University of Genoa, Italy</i>	
CHAL-3: MULTI-CHANNEL MULTI-PARTY MEETING TRANSCRIPTION	
CHAL-3.1: SUMMARY ON THE ICASSP 2022 MULTI-CHANNEL MULTI-PARTY MEETING TRANSCRIPTION GRAND CHALLENGE	9156
<i>Fan Yu, Shiliang Zhang, Zhihao Du, Siqi Zheng, Weilong Huang, Zhijie Yan, Bin Ma, Alibaba, China; Pengcheng Guo, Yihui Fu, Lei Xie, AISHELL Foundation, China; Zheng-Hua Tan, Aalborg University, Denmark; DeLiang Wang, The Ohio State University, United States of America; Yanmin Qian, Shanghai Jiao Tong University, China; Kong Aik Lee, Institute for Infocomm Research, A*STAR, Singapore; Xin Xu, Hui Bu, Beijing Shell Shell Technology Co., Ltd., China</i>	

CHAL-3.2: THE CUHK-TENCENT SPEAKER DIARIZATION SYSTEM FOR THE ICASSP 2022 MULTI-CHANNEL MULTI-PARTY MEETING TRANSCRIPTION CHALLENGE	9161
<i>Naijun Zheng, Xixin Wu, Lingwei Meng, Jiawen Kang, Helen Meng, The Chinese University of Hong Kong, China; Na Li, Chao Weng, Dan Su, Tencent AI Lab, China; Haibin Wu, National Taiwan University, China</i>	
CHAL-3.3: THE USTC-XIMALAYA SYSTEM FOR THE ICASSP 2022 MULTI-CHANNEL MULTI-PARTY MEETING TRANSCRIPTION (M2MET) CHALLENGE	9166
<i>Maokui He, Xiaoqi Zhang, Yuxuan Wang, Shutong Niu, Jun Du, University of Science and Technology of China, China; Xiang Lv, Weilin Zhou, Jingjing Yin, Yuhang Cao, Heng Lu, Ximalaya Inc., China; Chin-Hui Lee, Georgia Institute of Technology, United States of America</i>	
CHAL-3.4: CROSS-CHANNEL ATTENTION-BASED TARGET SPEAKER VOICE ACTIVITY DETECTION: EXPERIMENTAL RESULTS FOR M2MET CHALLENGE	9171
<i>Weiying Wang, Duke University, United States of America; Xiaoyi Qin, Ming Li, Duke Kunshan University, China</i>	
CHAL-3.5: THE VOLCSPEECH SYSTEM FOR THE ICASSP 2022 MULTI-CHANNEL MULTI-PARTY MEETING TRANSCRIPTION CHALLENGE	9176
<i>Chen Shen, Yi Liu, Wenzhi Fan, Bin Wang, Shixue Wen, Yao Tian, Jun Zhang, Jingsheng Yang, Zejun Ma, bytedance, China</i>	
CHAL-3.6: THE ROYALFLUSH SYSTEM OF SPEECH RECOGNITION FOR M2MET CHALLENGE	9181
<i>Shuaishuai Ye, Peiyao Wang, Shunfei Chen, Xinhui Hu, Xinkang Xu, Hithink RoyalFlush Information Network Co.,Ltd., China</i>	
 CHAL-4: L3DAS22 MACHINE LEARNING FOR 3D AUDIO SIGNAL PROCESSING	
CHAL-4.1: L3DAS22 CHALLENGE: LEARNING 3D AUDIO SOURCES IN A REAL OFFICE ENVIRONMENT	9186
<i>Eric Guizzo, Christian Marinoni, Marco Pennese, Aurelio Uncini, Danilo Comminiello, Sapienza University of Rome, Italy; Xinlei Ren, Xiguang Zheng, Chen Zhang, Kuaishou Technology, China; Bruno Masiero, University of Campinas, Brazil</i>	
CHAL-4.2: ICASSP 2022 L3DAS22 CHALLENGE: ENSEMBLE OF RESNET-CONFORMERS WITH AMBISONICS DATA AUGMENTATION FOR SOUND EVENT LOCALIZATION AND DETECTION	9191
<i>Yongjian Mao, Ying Zeng, Hongqing Liu, Wenbin Zhu, Yi Zhou, Chongqing University of Posts and Telecommunications, China</i>	
CHAL-4.3: A TRACK-WISE ENSEMBLE EVENT INDEPENDENT NETWORK FOR POLYPHONIC SOUND EVENT LOCALIZATION AND DETECTION	9196
<i>Jinbo Hu, Ming Wu, Feiran Yang, Jun Yang, Institute of Acoustics, Chinese Academy of Sciences, China; Yin Cao, Mark Plumbley, University of Surrey, United Kingdom of Great Britain and Northern Ireland; Qiuqiang Kong, ByteDance, China</i>	
CHAL-4.4: TOWARDS LOW-DISTORTION MULTI-CHANNEL SPEECH ENHANCEMENT: THE ESPNET-SE SUBMISSION TO THE L3DAS22 CHALLENGE	9201
<i>Yen-Ju Lu, Xuankai Chang, Zhong-Qiu Wang, Shinji Watanabe, Carnegie Mellon University, United States of America; Samuele Cornell, Università Politecnica delle Marche, Italy; Wangyou Zhang, Chenda Li, Shanghai Jiao Tong University, China; Zhaozheng Ni, Meta, United States of America</i>	
CHAL-4.5: MULTI-SCALE TEMPORAL FREQUENCY CONVOLUTIONAL NETWORK WITH AXIAL ATTENTION FOR MULTI-CHANNEL SPEECH ENHANCEMENT	9206
<i>Guochang Zhang, Chunliang Wang, Libiao Yu, Jianqiang Wei, Baidu, China</i>	
CHAL-4.6: THE PCG-AIID SYSTEM FOR L3DAS22 CHALLENGE: MIMO AND MISO CONVOLUTIONAL RECURRENT NETWORK FOR MULTI CHANNEL SPEECH ENHANCEMENT AND SPEECH RECOGNITION	9211
<i>Jingdong Li, Yuanyuan Zhu, Dawei Luo, Yun Liu, Guohui Cui, Zhaoxia Li, Tencent, China</i>	

CHAL-5: AUDIO DEEPFAKE DETECTION

CHAL-5.1: ADD 2022: THE FIRST AUDIO DEEP SYNTHESIS DETECTION 9216 CHALLENGE

Jiangyan Yi, Ruibo Fu, Jianhua Tao, Shuai Nie, Haoxin Ma, Chenglong Wang, Tao Wang, Zhengkun Tian, Ye Bai, Cunhang Fan, Shan Liang, Shiming Wang, Shuai Zhang, Zhengqi Wen, Institute of Automation, Chinese Academy of Sciences, China; Xinrui Yan, Le Xu, University of Chinese Academy of Sciences, China; Haizhou Li, National University of Singapore, Singapore

CHAL-5.2: TIME DOMAIN ADVERSARIAL VOICE CONVERSION FOR ADD 2022..... 9221

Cheng Wen, Tingwei Guo, Xingjun Tan, Rui Yan, Shuran Zhou, Chuandong Xie, Wei Zou, Xiangang Li, Beike, China

CHAL-5.3: AUDIO DEEPFAKE DETECTION SYSTEM WITH NEURAL STITCHING 9226 FOR ADD 2022

Rui Yan, Cheng Wen, Shuran Zhou, Tingwei Guo, Wei Zou, Xiangang Li, Beike, China

CHAL-5.4: FAKE AUDIO DETECTION BASED ON UNSUPERVISED PRETRAINING 9231 MODELS

Zhiqiang Lv, Shanshan Zhang, Kai Tang, Pengfei Hu, Tencent, China

CHAL-5.5: PARTIALLY FAKE AUDIO DETECTION BY SELF-ATTENTION-BASED FAKE 9236 SPAN DISCOVERY

Haibin Wu, Hung-yi Lee, National Taiwan University, Taiwan; Heng-Cheng Kuo, Kuo-Hsuan Hung, Yu Tsao, Research Center for Information Technology Innovation, Academia Sinica, Taiwan, Taiwan; Naijun Zheng, Helen Meng, The Chinese University of Hong Kong, Hong Kong; Hsin-Min Wang, Institute of Information Science, Academia Sinica, Taiwan, Taiwan

CHAL-5.6: THE VICOMTECH AUDIO DEEPFAKE DETECTION SYSTEM BASED ON 9241 WAV2VEC2 FOR THE 2022 ADD CHALLENGE

Juan M. Martín-Doñas, Aitor Álvarez, Vicomtech Foundation, Spain

CHAL-6: MULTIMODAL INFORMATION BASED SPEECH PROCESSING

CHAL-6.1: AUDIO-VISUAL WAKE WORD SPOTTING SYSTEM FOR MISP 9246 CHALLENGE 2021

Yanguang Xu, Jianwei Sun, Yang Han, Shuaijiang Zhao, Chaoyang Mei, Tingwei Guo, Shuran Zhou, Chuandong Xie, Wei Zou, Xiangang Li, Beike, China

CHAL-6.2: CHANNEL-WISE AV-FUSION ATTENTION FOR MULTI-CHANNEL 9251 AUDIO-VISUAL SPEECH RECOGNITION

Gaopeng Xu, Song Yang, Wei Li, Sang Wang, Wei Guo, Junfeng Yuan, Jie Gao, NIO Co., Ltd., China

CHAL-6.3: THE DKU AUDIO-VISUAL WAKE WORD SPOTTING SYSTEM FOR THE 9256 2021 MISP CHALLENGE

Ming Cheng, Haoxu Wang, Ming Li, Wuhan University, China; Yechen Wang, Duke Kunshan University, China

CHAL-6.4: THE SJTU SYSTEM FOR MULTIMODAL INFORMATION BASED SPEECH 9261 PROCESSING CHALLENGE 2021

Wei Wang, Xun Gong, Yifei Wu, Zhikai Zhou, Chenda Li, Wangyou Zhang, Bing Han, Yanmin Qian, Shanghai Jiao Tong University, China

CHAL-6.5: THE FIRST MULTIMODAL INFORMATION BASED SPEECH PROCESSING 9266 (MISP) CHALLENGE: DATA, TASKS, BASELINES AND RESULTS

Hang Chen, Hengshun Zhou, Jun Du, University of Science and Technology of China, China; Chin-Hui Lee, Georgia Institute of Technology, United States of America; Jingdong Chen, Northwestern Polytechnical University, China; Shinji Watanabe, Carnegie Mellon University, United States of America; Sabato Marco Siniscalchi, Kore University of Enna, Italy; Odette Scharenborg, Delft University of Technology, Netherlands; Di-Yuan Liu, Bao-Cai Yin, Jia Pan, Jian-Qing Gao, Cong Liu, iFlytek, China

CHAL-7: DEEP NOISE SUPPRESSION

CHAL-7.1: ICASSP 2022 DEEP NOISE SUPPRESSION CHALLENGE 9271
Harishchandra Dubey, Vishak Gopal, Ross Cutler, Ashkan Aazami, Sergiy Matuskevych, Sebastian Braun, Sefik Emre Eskimez, Manthan Thakker, Takuya Yoshioka, Hannes Gamper, Robert Aichner, Microsoft Corporation, Redmond, WA, USA, United States of America

CHAL-7.2: FB-MSTCN: A FULL-BAND SINGLE-CHANNEL SPEECH ENHANCEMENT 9276
METHOD BASED ON MULTI-SCALE TEMPORAL CONVOLUTIONAL NETWORK
Zehua Zhang, Lu Zhang, Xuyi Zhuang, Yukun Qian, Heng Li, Mingjiang Wang, Harbin Institute of Technology(Shenzhen), China

CHAL-7.3: FRCRN: BOOSTING FEATURE REPRESENTATION USING FREQUENCY 9281
RECURRENCE FOR MONAURAL SPEECH ENHANCEMENT
Shengkui Zhao, Bin Ma, Alibaba Group, Singapore; Karn N. Watcharasupat, Woon-Seng Gan, Nanyang Technological University (NTU), Singapore

CHAL-7.4: HARMONIC GATED COMPENSATION NETWORK PLUS FOR ICASSP 2022 9286
DNS CHALLENGE
Tianrui Wang, Weibin Zhu, Beijing Jiaotong University, China; Yingying Gao, Yanan Chen, Junlan Feng, Shilei Zhang, China Mobile Research Institute, China

CHAL-7.5: TEA-PSE: TENCENT-ETHEREAL-AUDIO-LAB PERSONALIZED SPEECH 9291
ENHANCEMENT SYSTEM FOR ICASSP 2022 DNS CHALLENGE
Yukai Ju, Xiaopeng Yan, Yihui Fu, Shubo Lv, Luyao Cheng, Lei Xie, Northwestern Polytechnical University, China; Wei Rao, Yannan Wang, Shidong Shang, Tencent, China

CHAL-7.6: MULTI-STAGE AND MULTI-LOSS TRAINING FOR FULLBAND 9296
NON-PERSONALIZED AND PERSONALIZED SPEECH ENHANCEMENT
Lianwu Chen, Chenglin Xu, Xu Zhang, Xinlei Ren, Xiguang Zheng, Chen Zhang, Liang Guo, Bing Yu, Kuaishou Technology, China

CHAL-8: ROOT CAUSE ANALYSIS FOR WIRELESS NETWORK FAULTS LOCALIZATION

CHAL-8.1: ICASSP-SPGC 2022: ROOT CAUSE ANALYSIS FOR WIRELESS NETWORK 9301
FAULT LOCALIZATION
Tianjian Zhang, Qian Chen, Feng Yin, Zhi-Quan Luo, The Chinese University of Hong Kong, Shenzhen; Shenzhen Research Institute of Big Data, China; Yi Jiang, The Chinese University of Hong Kong, Shenzhen, China; Dandan Miao, Tao Quan, Huawei Technologies Co., Ltd, China; Qingjiang Shi, Shenzhen Research Institute of Big Data, China

CHAL-8.2: ACCURATE INFERENCE OF UNSEEN COMBINATIONS OF MULTIPLE 9306
ROOTCAUSES WITH CLASSIFIER ENSEMBLE
Xuan Zhang, Longxiang Xiong, Ningyuan Sun, Mingxia Wang, Hao Tang, Network Intelligent Laboratory, School of Electronics and Information Engineering, Beijing Jiaotong University, China; Yanxing Zhao, Beijing Baolande Software Corporation, China

CHAL-8.3: CAUSAL ALIGNMENT BASED FAULT ROOT CAUSES LOCALIZATION FOR 9311
WIRELESS NETWORK
Yuequn Liu, Wenhui Zhu, Jie Qiao, Zhiyi Huang, Yu Xiang, Xuanchi Chen, Wei Chen, Ruichu Cai, Guangdong University of Technology, China

CHAL-8.4: NETRCA: AN EFFECTIVE NETWORK FAULT CAUSE LOCALIZATION 9316
ALGORITHM
Chaoli Zhang, Zhiqiang Zhou, Yingying Zhang, Linxiao Yang, Kai He, Qingsong Wen, Liang Sun, Alibaba Group, China