

Daily Schedule

Tutorials

INTERSPEECH 2018 - September 2, 2018 - PROGRAM AT A GLANCE						
	MR G.01-G.03	MR G.04-G.06	MR 1.03	MR 1.04		
08:00-08:30						REGISTRATION at Ground Floor
08:30-09:00						
09:00-09:30	Tutorial 1: Deep Learning Based Speech Separation	Tutorial 2: End-To-End Models for Automatic Speech Recognition	Tutorial 3: Spoofing Attacks in Automatic Speaker Verification: Analysis and Countermeasures	Tutorial 4: Spoken Dialog Technology for Education Domain Applications		
09:30-10:00						
10:00-10:30						
10:30-11:00						
	Refreshment Break (First Floor Foyer)					
11:00-11:30	Tutorial 1: Deep Learning Based Speech Separation	Tutorial 2: End-To-End Models for Automatic Speech Recognition	Tutorial 3: Spoofing Attacks in Automatic Speaker Verification: Analysis and Countermeasures	Tutorial 4: Spoken Dialog Technology for Education Domain Applications		
11:30-12:00						
12:00-12:30						
12:30-13:00	Lunch Break					
13:00-13:30						
13:30-14:00						
14:00-14:30	Tutorial 5: Information Theory of Deep Learning: What do the Layers of Deep Neural Networks Represent ?	Tutorial 6: Articulatory Representations: Measurement, Estimation, and Application to State-of-the-Art Automatic Speech Recognition	Tutorial 7: Vocal Music Processing for Music Information Retrieval (MIR)	Tutorial 8: Multimodal Speech and Audio Processing in Audio-Visual Human-Robot Interaction		
14:30-15:00						
15:00-15:30						
15:30-16:00	Refreshment Break (First Floor Foyer)					
16:00-16:30	Tutorial 5: Information Theory of Deep Learning: What do the Layers of Deep Neural Networks Represent ?	Tutorial 6: Articulatory Representations: Measurement, Estimation, and Application to State-of-the-Art Automatic Speech Recognition	Tutorial 7: Vocal Music Processing for Music Information Retrieval (MIR)	Tutorial 8: Multimodal Speech and Audio Processing in Audio-Visual Human-Robot Interaction		
1630-17:00						
17:00-17:30						
17:30-18:00						

Day 1 (Monday)

INTERSPEECH 2018 - September 3, 2018 - PROGRAM AT A GLANCE										
	Hall 1	Hall 2	Hall 3	MR G.01-G.02	MR G.03-G.04	MR G.05-G.06	MR 1.01-1.02	Hall 4 - 6		
08:00-08:30										
08:30-09:00										
09:00-09:30	Welcome Refreshment (Hall 5-6)								REGISTRATION at Ground Floor	
09:30-10:00										
10:00-10:30			Opening Ceremony							
10:30-11:00										
11:00-11:30			ISCA Medal Talk: Bishnu S. Atal							
11:30-12:00										
12:00-12:30	Lunch Break									
12:30-13:00										
13:00-13:30										
13:30-14:00										
14:00-14:30										
14:30-15:00	O-1-2: Prosody Modeling and Generation	O-1-3: Speaker Verification I	O-1-1: End-to-End Speech Recognition	O-1-4: Spoken Term Detection		S&T-1-1: Show and Tell 1	SS-1-1: The INTERSPEECH 2018 Computational Paralinguistics ChallengeE	Posters: P-1-1 P-1-2 P-1-3 P-1-4	Speaker Check-in (MR 1.07)	
15:00-15:30										
15:30-16:00										
16:00-16:30	Refreshment Break (Hall 5-6)									
1630-17:00	O-2-1: ASR Systems and Technologies	O-2-2: Deception, Personality, and Culture Attribute		O-2-3: Automatic Detection and Recognition of Voice and Speech Disorders	O-2-4: Voice Conversion	S&T-2-1: Show and Tell 2	SS-2-1: The INTERSPEECH 2018 Computational Paralinguistics ChallengeE	Posters: P-2-1 P-2-2 P-2-3 P-2-4 P-2-5		
17:00-17:30										
17:30-18:00										
18:00-18:30										
18:30-19:00			Cultural Programme							
19:00-19:30										
19:30-20:00	Welcome Reception (Novotel Gardens)									
20:00-20:30										
20:30-21:00										
21:00-21:30										
P-1-1: Speech Segments and Voice Quality P-1-2: Speaker State and Trait P-1-3: Deep Learning for Source Separation and Pitch Tracking P-1-4: Acoustic Analysis-Synthesis of Speech Disorders				P-2-1: Spoken Dialogue Systems and Conversational Analysis P-2-2: Spoofing Detection P-2-3: Speech Analysis and Representation P-2-4: Sequence Models for ASR P-2-5: Source Separation and Spatial Analysis						

Day 2 (Tuesday)

INTERSPEECH 2018 - September 4, 2018 - PROGRAM AT A GLANCE								
	Hall 1	Hall 2	Hall 3	MR G.01-G.02	MR G.03-G.04	MR G.05-G.06	MR 1.01-1.02	Hall 4 - 6
08:00-08:30								
08:30-09:00			Plenary Talk-1: Jacqueline Vaissière					
09:00-09:30								
09:30-10:00	Refreshment Break (Hall 4-5)							
10:00-10:30	O-1-2: Statistical Parametric Speech Synthesis	O-1-3: Emotion Modeling	O-1-1: Acoustic Model Adaptation	O-1-4: Models of Speech Perception	O-1-5: Multimodal Dialogue Systems	S&T-1-1: Show and Tell 3	SS-1-1: Speech Recognition for Indian Languages	Posters: P-1-1 P-1-2 P-1-3 P-1-4
10:30-11:00								
11:00-11:30								
11:30-12:00								
12:00-12:30			Perspective Talk-1: Dilek Hakkani-Tur					
12:30-13:00	Industry Presentation- 2: JD	Industry Presentation- 3: Uniphore	Industry Presentation- 1: Amazon					
13:00-13:30	Lunch Break							
13:30-14:00								
14:00-14:30								
14:30-15:00	O-2-2: Audio Events and Acoustic Scenes	O-2-3: Speaker Diarization	O-2-1: Dereverberat ion	O-2-4: Phonation	O-2-5: Cognition and Brain Studies	S&T-2-1: Show and Tell 4	SS-2-1: Deep Neural Networks: How Can We Interpret What They Learned?	Posters: P-2-1 P-2-2 P-2-3 SS-2-2 P-2-5
15:00-15:30								
15:30-16:00								
16:00-16:30								
16:30-17:00	Refreshment Break (Hall 5-6)							
17:00-17:30			Perspective Talk-2: Bhuvana Ramabhadra n					
17:30-18:00			ISCA General Assembly					
18:00-18:30								
18:30-19:00								
19:00-19:30	Students' Reception (Club House, Jesmin, Ella Hotel, Gachibowli) Reviewers' Reception (Kebab Pavilion, Ella Hotel, Gachibowli)							
19:30-20:00								
20:00-20:30								
20:30-21:00								
21:00-21:30								
P-1-1: Speaker Verification II P-1-2: Novel Approaches to Enhancement P-1-3: Syllabification, Rhythm, and Voice Activity Detection P-1-4: Selected Topics in Neural Speech Processing					P-2-1: Speech and Singing Production P-2-2: Robust Speech Recognition P-2-3: Applications in Education and Learning SS-2-2: Integrating Speech Science and Technology for Clinical Applications P-2-5: Speaker Characterization and Analysis			

REGISTRATION at Ground Floor

Speaker Check-in (MR 1.07)

Day 3 (Wednesday)

INTERSPEECH 2018 - September 5, 2018 - PROGRAM AT A GLANCE										
	Hall 1	Hall 2	Hall 3	MR G.01-G.02	MR G.03-G.04	MR G.05-G.06	MR 1.01-1.02	Hall 4 - 6		
08:00-08:30										REGISTRATION at Ground Floor Speaker Check-in (MR 1.07)
08:30-09:00			Plenary Talk-2: Hervé Bourlard							
09:00-09:30										
09:30-10:00	Refreshment Break (Hall 5 - 6)									
10:00-10:30	O-1-2: Language Identification	O-1-3: Production of Prosody	O-1-1: Novel Neural Network Architectures for Acoustic Modelling	O-1-4: Speech Intelligibility and Quality	SS-1-1: Integrating Speech Science and Technology for Clinical Applications	S&T-1-1: Show and Tell 5	SS-1-2: Speech Technologies for Code-Switching in Multilingual Communities	Posters: P-1-1 P-1-2 P-1-3 P-1-4		
10:30-11:00										
11:00-11:30										
11:30-12:00										
12:00-12:30			Perspective Talk-3: Nima Mesgarani							
12:30-13:00	Industry Presentation-5: Xiaomi	Industry Presentation -6: MeitY	Industry Presentation-4: Microsoft							
13:00-13:30	Lunch Break									
13:30-14:00										
14:00-14:30										
14:30-15:00	O-2-2: Speaker Verification Using Neural Network Methods I	O-2-3: Speech Perception in Adverse Conditions	O-2-1: Recurrent Neural Models for ASR	O-2-4: Measuring Pitch and Articulation	O-2-5: Speech and Language Analytics for Mental Health	S&T-1-2: Show and Tell 6	SS-2-1: Spoken CALL Shared Task, Second Edition	Posters: P-2-1 P-2-2 P-2-3 P-2-4		
15:00-15:30										
15:30-16:00										
16:00-16:30										
16:30-17:00	Refreshment Break (Hall 5 - 6)									
17:00-17:30	O-3-1: Zero-resource Speech Recognition	O-3-2: Spatial and Phase Cues for Source Separation and Speech Recognition		O-3-3: Dialectal Variation	O-3-4: Spoken Corpora and Annotation		SS-3-1: The First DIHARD Speech Diarization Challenge	Posters: P-3-1 P-3-2 P-3-3 P-3-4		
17:30-18:00										
18:00-18:30										
18:30-19:00										
19:00-19:30	Cultural Programme + Banquet (Hall 5 - 6)									
19:30-20:00										
20:00-20:30										
20:30-21:00										
21:00-21:30										
P-1-1: Voice Conversion and Speech Synthesis P-1-2: Extracting Information from Audio P-1-3: Signal Analysis for the Natural, Biological and Social Sciences P-1-4: Speech Prosody			P-2-1: Adjusting to Speaker, Accent, and Domain P-2-2: Speech Synthesis Paradigms and Methods P-2-3: Second Language Acquisition and Code-switching P-2-4: Topics in Speech Recognition			P-3-1: Text Analysis, Multilingual Issues and Evaluation in Speech Synthesis P-3-2: Neural Network Training Strategies for ASR P-3-3: Application of ASR in Medical Practice P-3-4: Source and Supra-segmentals				

Day 4 (Thursday)

INTERSPEECH 2018 - September 6, 2018 - PROGRAM AT A GLANCE								
	Hall 1	Hall 2	Hall 3	MR G.01-G.02	MR G.03-G.04	MR G.05-G.06	MR 1.01-1.02	Hall 4 - 6
08:00-08:30								
08:30-09:00			Plenary Talk-3: Helen Meng					
09:00-09:30								
09:30-10:00	Refreshment Break (Hall 5 - 6)							
10:00-10:30	O-1-2: Expressive Speech Synthesis	O-1-3: Representation Learning for Emotion	O-1-1: Distant ASR	O-1-4: Articulatory Information, Modeling and Inversion	SS-1-1: Novel Paradigms for Direct Synthesis Based on Speech-Related Biosignals	S&T-1-1: Show and Tell 7	SS-1-2: Low Resource Speech Recognition Challenge for Indian Languages	Posters: P-1-1 P-1-2 P-1-3 P-1-4
10:30-11:00								
11:00-11:30								
11:30-12:00								
12:00-12:30			Perspective Talk-4: Sriram Ganapathy					
12:30:13:00	Industry Presentation- 8: Baidu	Industry Presentation- 9: Nvidia	Industry Presentation- 7: Samsung					
13:00-13:30	Lunch Break							
13:30-14:00								
14:00-14:30								
14:30-15:00	O-2-2: Source Separation from Monaural Input	O-2-3: Multimodal Systems	O-2-1: Spoken Language Understanding	O-2-4: Coding				Posters: P-2-1 P-2-2 P-2-3 P-2-4
15:00-15:30								
15:30-16:00								
16:00-16:30								
1630-17:00	Refreshment Break (Hall 5 - 6)							
17:00-17:30			Closing Ceremony					
17:30-18:00								
P-1-1: Deep Enhancement P-1-2: Acoustic Scenes and Rare Events P-1-3: Language Modeling P-1-4: Speech Pathology, Depression, and Medical Applications					P-2-1: Speaker Verification Using Neural Network Methods II P-2-2: Emotion Recognition and Analysis P-2-3: Acoustic Modelling P-2-4: Speech and Speaker Perception			