

Midwest Voice

Abstract Booklet

November 1, 2025
Voxman Music Building
University of Iowa
Iowa City, IA

Stay tuned for future opportunities and events



Titles and Authors

KEYNOTE “TO METHOD OR NOT TO METHOD: THAT IS THE QUESTION”	4
<i>Kittie Verdolini Abbott, PhD, CCC-SLP, PAVA-RV, MDiv</i>	4
SUBCORTICAL CONTRIBUTIONS TO VOICE AND SPEECH PRODUCTION: INSIGHTS FROM DIRECT HUMAN BRAIN RECORDINGS	5
<i>Andrea H. Rohl¹, Joel I. Berger¹, Karim Johari², Jeremy D.W. Greenlee^{1,3}</i>	5
QUANTIFYING TALK-TIME IN MUSIC INSTRUCTION: AN AUTOMATED METHOD	6
<i>Shanshan Zhang, MM¹, Christian T. Herbst, MA, PhD², Katrina Miller, Ed.D., CCC-SLP³, David Meyer, DM^{4, 5, 6}</i>	6
THE WIRED VOICE: INTEGRATING TECHNOLOGY INTO CONTEMPORARY VOCAL PRACTICE	7
<i>Matt Edwards</i>	7
SINGING VOICE MACHINE LEARNING AI MODELS AS DIAGNOSTIC BIOMARKERS FOR VOCAL HEALTH, DYSPHONIA, TMD, AND TREMOR	8
<i>Theodora Nestorova^{1,2}, Jiarui Xie³, Chonghui Zhang³, Felicia Nyanyo³, Yaoyao Fiona Zhao^{3,4}, Luc Mongeau⁴</i>	8
SIX-PACK SINGING: ON ASSOCIATIONS BETWEEN ANTERIOR ABDOMINAL MUSCLE ACTIVITY AND VITAL CAPACITY PERCENTAGE	10
<i>Christian T. Herbst^{1,2}, Anna Miles³, Jacqueline Allen^{3,4}, Calvin P. Baker^{3,5}</i>	10
KEYNOTE Sounding the Self: Vocality in Performance, Pedagogy, and Care	11
<i>Alexis Davis-Hazell^{1,2}</i>	11
The Valsalva Ear Pressure Equalization Maneuver in Operatic Singers: A Qualitative Pilot Study	12
<i>Ian DeNolfo, MM¹, Gioacchino LiVigni^{2,3}, David Meyer, DM^{4,5,6}, Robert T. Sataloff, MD, DMA, FACS^{7,8}</i>	12
TOWARD EVIDENCE-BASED CHORAL CONDUCTING GESTURE	13
<i>Jeremy N. Manternach</i>	13
TREMBLE CLEFS: THERAPEUTIC GROUP SINGING FOR INDIVIDUALS WITH PARKINSON’S DISEASE	14
<i>Dr. Abbey L. Dvorak¹, Dr. Sun Joo Lee²</i>	14
KEYNOTE Voice disorders that present without dysphonia	15
<i>Thomas L. Carroll, MD</i>	15
DECODING ILLUSORY MELODY FROM ELECTROENCEPHALOGRAPHY USING BACKWARD TEMPORAL RESPONSE FUNCTIONS	16
<i>Ahyeon Choi and Inyong Choi</i>	16

RETHINKING AUDITORY-PERCEPTUAL SCALING: ADVANCING RESEARCH AND CLINICAL IMPACT IN SPEECH AND VOICE ASSESSMENT.....	17
<i>Mili Kuruvilla-Dugdale.....</i>	<i>17</i>
AI-GENERATED DYSPHONIC SPEECH: MORE CLEAR, LESS REAL?	18
<i>Pasquale Bottalico¹, Charles J. Nudelman^{1,2}, Daniel Fogerty¹, Virginia Tardini^{1,3}, Keiko Ishikawa⁴</i>	<i>18</i>
ASSESSMENT OF VOCALIZATION IN A RAT AND FINDINGS FROM THE <i>PINK1</i>-/- RAT MODEL OF PARKINSON DISEASE.....	19
<i>Maryann N. Krasko, PhD^{1,2,3,4}, Denis Michael Rudisch, MM, Dipl Univ^{1,2,5}, Michelle R. Ciucci, PhD, CCC-SLP^{1,2,6}</i>	<i>19</i>
UNISONO AFTER LARYNGECTOMY: FINDING A NEW VOICE THROUGH MUSIC AND SCIENCE	21
<i>Thomas Moors^{1,2}, David Howard³, Evangelos Himonides⁴, David Meyer^{5,6,7}</i>	<i>21</i>
KEYNOTE THE BIOMECHANICS OF VOCAL FOLD TRAUMA	22
<i>Jack J. Jiang, MD, PhD</i>	<i>22</i>
EARLY DETECTION OF VOICE SYMPTOMS IN PRE-SERVICE TEACHERS OF SOUTH AMERICA.....	23
<i>Lady Catherine Cantor-Cutiva¹, Léslie Piccolotto Ferreira², Maria Lúcia Vaz Masson³, Iara Bittante de Oliveira⁴, Maria Celina Malebran Bezerra de Mello⁵, Alejandro Morales Quevedo⁶, Carlos Manzano⁷, María del Carmen Dalmasso⁸, Adriana Díaz Gutiérrez⁹, Jessica Laura Ramonda¹⁰, Eric J. Hunter¹¹</i>	<i>23</i>
MULTI-MODAL MRI AND CT IMAGE PROCESSING FOR VOCAL TRACT AREA FUNCTION ESTIMATION	25
<i>Swati Ramtilak ¹, David Meyer ^{2,5,6}, Aiming Lu ⁴, James Holmes ^{3,1}, Jarron Atha ³, Eric Hoffman ^{3,1}, Sajan Goud Lingala ^{1,3}</i>	<i>25</i>
A 3D UPPER AIRWAY ATLAS: DATA, SEGMENTATIONS, AND PRINTABLE VOCAL TRACTS	26
<i>Subin Erattakulangara^{1,2}, Swati Ramtilak¹, Karthika Kelat¹, David Meyer^{3,4,5}, Sajan Goud Lingala^{1,6}</i>	<i>26</i>
TRANSFORMING VOCAL VARIABILITY FROM NOISE TO SIGNAL: A PRECISION MEDICINE APPROACH TO VOICE HEALTH	28
<i>Eric J. Hunter, Mark L. Berardi.....</i>	<i>28</i>
FROM SIGNAL VARIABILITY TO VOCAL SUSTAINABILITY: OPERATIONALIZING PRECISION MEDICINE FOR OCCUPATIONAL VOICE HEALTH.....	29
<i>Mark L. Berardi, Eric J. Hunter.....</i>	<i>29</i>

KEYNOTE**“TO METHOD OR NOT TO METHOD: THAT IS THE QUESTION”****KITTIE VERDOLINI ABBOTT, PhD, CCC-SLP, PAVA-RV, MDiv**

Department Communication Sciences and Disorders, University of Delaware, Newark, Delaware, USA

ABSTRACT

In voice pedagogy, most teachers have an “angle.” For example, some teachers may be all about voice registers. Other teachers may be all about breathing. Others still may be about vocal tract manipulations. “Favorite exercises” are used as methods to address the “angle’s end.” In contrast, some teachers may utilize a more formalized, structured training method. The Estill approach comes to mind. Either way, the “method” approach – whether grounded in “angles” or a formal structure -- has many advantages, not least of which is the codification of teaching honed and refined over time. Lack of “method” often means inventing lessons in the moment, moment to moment, in a way that can easily leave the teacher lost if not panicked (I say from experience). In this presentation, I will talk about a “third way.” In this approach, we indeed create lessons in the moment, responding to the student we have before us in that day and that hour, but with a twist. The twist is that although each session may be new in content, the content is generated by a set of relatively constant “rules” emerging from the science of perceptual-motor learning. In this presentation, I will discuss some of the major findings in contemporary motor learning literature and illustrate how they may be flexibly applied in teaching with the performing artist we find before us, in a way that is vital, alive, and grounded in principle.

KEYWORDS: Voice Pedagogy, Perceptual-Motor Learning

PRESENTER BIO: Kittie Verdolini Abbott, PhD, CCC-SLP, MDiv, is a Professor of Communication Sciences and Disorders and Linguistics and Cognitive Science at the University of Delaware. Her research focuses on voice and voice disorders, with interests spanning hydration, wound healing, emotions, cognition, motor learning, exercise physiology, and clinical trials. She is a singer, teacher of singing, a Lessac-trained actor, and an ordained American Baptist minister.

ACKNOWLEDGMENTS: Funding from the National Institute on Deafness and Other Communication Disorders is acknowledged.

DISCLOSURES: Dr. Verdolini Abbott receives honoraria from Visions in Voice, LLC, for continuing education seminars, and royalties for publications from Plural Publishing Co. and the National Center for Voice and Speech. She is a salaried professor at the University of Delaware and her research is funded by the National Institutes of Health.

REFERENCES

Schmidt, R.A., Lee, T.D., Winstein, C., Wulf, G., & Zelaznik, H.N. (2019). *Motor Control and Learning: A Behavioral Emphasis* (6th Edition). Champaign, IL: Human Kinetics.

Titze, I.R. & Verdolini Abbott, K. (2012). *Vocology*. Iowa City, Iowa: National Center for Voice and Speech.

SUBCORTICAL CONTRIBUTIONS TO VOICE AND SPEECH PRODUCTION: INSIGHTS FROM DIRECT HUMAN BRAIN RECORDINGS

ANDREA H. ROHL¹, JOEL I. BERGER¹, KARIM JOHARI², JEREMY D.W. GREENLEE^{1,3}

¹ Department of Neurosurgery, University of Iowa, Iowa City, IA USA

² Department of Communication Sciences and Disorders, Louisiana State University, Baton Rouge, LA USA

³ Department of Otolaryngology, University of Iowa, Iowa City, IA USA

ABSTRACT

While the cortical networks supporting speech production have been extensively studied using functional magnetic resonance imaging, lesion-symptom mapping, and electrocorticography, the contribution of subcortical structures remains less well-understood. These regions are more challenging to study due to their anatomical depth and proximity to adjacent structures. As a result, current models of speech production often underrepresent the role of the basal ganglia and thalamus, despite their known importance for motor control and cognitive integration.

Awake deep brain stimulation implantation surgeries provide a unique opportunity to study human subcortical function during naturalistic behavior. Using neural recordings collected intraoperatively during speech and motor tasks, we examined single-neuron and local field potential (LFP) activity within the subthalamic nucleus (STN) and motor thalamus in individuals with Parkinson's disease (PD) and essential tremor respectively. Participants performed speech production tasks varying in linguistic complexity (sustained vowels, syllable repetition, and sentence production) as well as orofacial motor tasks designed to isolate articulator movement from syntactic and semantic contexts.

In the motor thalamus, LFP data show significant increases in low frequency activity with increasing task complexity. This may indicate greater need for long-range cortico-thalamic communication for speech tasks of greater complexity. From single neuron recordings in the STN of PD patients we found heterogeneous activity patterns, with subsets of neurons responding to different speech tasks indicating potential task-specific neuronal organization within the STN. A greater proportion of STN neurons showed changes in firing rate with speech tasks compared to simple finger tapping tasks, which may reflect greater complexity of speech tasks.

Our findings demonstrate distinct neuronal modulation patterns across subcortical nuclei and task types. Both regions exhibited activity tightly time-locked to speech onset, suggesting a role in the initiation of speech motor programs. In both the STN and the motor thalamus, task-dependent modulation varied with linguistic complexity, indicating sensitivity to higher-order aspects of speech production beyond pure motor execution.

These results provide converging evidence that subcortical nuclei contribute differentially to speech production—supporting both motor and cognitive components of speech. By linking single-neuron activity to specific speech behaviors, this work extends models of speech control to include the basal ganglia–thalamocortical network as an integral substrate for the coordination and initiation of complex speech actions. Understanding how these subcortical circuits support speech will both advance fundamental models of speech production and inform the development of targeted neuromodulatory therapies to preserve or restore communication in neurodegenerative diseases.

KEYWORDS: Neurophysiology, Speech, Parkinson's Disease, Deep Brain Stimulation

PRESENTER BIO: Annie Rohl is a speech-language pathologist and Neuroscience PhD candidate at the University of Iowa. Her research focuses on the impact of medical and surgical treatments for movement disorders on speech production. The overarching goal of her research is to optimize life participation in patients with neurodegenerative disease.

ACKNOWLEDGMENTS: This work was supported by a grant from the NIDCD (R01DC017718; PI: Jeremy Greenlee).

DISCLOSURES: None of the authors have conflicts of interest to disclose.

REFERENCES

Johari, K., Berger, J., Rohl, A., & Greenlee, J. (in press). Dissociation between simple and complex speech motor tasks within bilateral motor thalamus. *eNeuro*.

QUANTIFYING TALK-TIME IN MUSIC INSTRUCTION: AN AUTOMATED METHOD

SHANSHAN ZHANG, MM¹, CHRISTIAN T. HERBST, MA, PhD², KATRINA MILLER, Ed.D., CCC-SLP³, DAVID MEYER, DM^{4,5,6}

¹ Voice Pedagogy, Shenandoah University, Winchester, VA, USA

² Department of Behavioural and Cognitive Biology, University of Vienna, Austria

³ Department of Speech Pathology, Shenandoah University, Winchester, VA, USA

⁴ University of Iowa School of Music, The University of Iowa, Iowa City, Iowa, USA

⁵ Department of Communication Sciences and Disorders, The University of Iowa, Iowa City, Iowa, USA

⁶ Shenandoah Conservatory, Winchester, Virginia, USA

ABSTRACT

Modern music pedagogy provides evidence-informed tools to help teachers optimize instruction. Preliminary studies have revealed that in many lessons, teachers spend more time talking than students spend singing or speaking—highlighting a common example of suboptimal teacher-student communication. Notably, current investigations of teachers' talk time in one-on-one lessons have all been conducted by manually analyzing the talk-time. This imposes the following limitations onto previous research: Manual video coding and analysis is both error-prone and time-intensive, resulting in limited sample size and thus weak statistical power. To address these challenges, this study introduces an automated method. Using Zoom's automatic transcription feature combined with a custom Python script, the study automates the coding and analysis of lesson recordings, making it possible to investigate teacher talk time on a larger scale. By offering a scalable, low-cost solution, this study aims to improve the efficiency and accessibility of research in music pedagogy.

KEYWORDS: Voice; Talk-time; Automation; Pedagogy

PRESENTER BIO: Shanshan Zhang is a DMA candidate in Voice Pedagogy at Shenandoah University, studying with Dr. David Meyer. Her research explores teacher talk time in one-on-one lessons. Active as both performer and educator in the DMV area, she brings passion for vocal artistry and pedagogy to students and audiences.

DISCLOSURES: None of the authors have conflicts of interest to disclose.

THE WIRED VOICE: INTEGRATING TECHNOLOGY INTO CONTEMPORARY VOCAL PRACTICE

MATT EDWARDS

Associate Professor of Voice and Theatre, Shenandoah University, Winchester, VA, United States

ABSTRACT

In Commercial and Musical Theatre singing, amplification is essential, but the microphone is rarely treated as an integral part of the vocal instrument. This session challenges that oversight by exploring how understanding microphone technology can transform vocal pedagogy and performance outcomes.

Participants will discover how the microphone's frequency response curve interacts with the performer's vocal output, and how this knowledge can be used to shape the singer's tone and reduce vocal fatigue. By learning to treat the microphone as a technical ally rather than a passive tool, voice teachers can help students improve their stamina while delivering more stylistically appropriate performances.

KEYWORDS: Singer tone, microphone feedback

PRESENTER BIO: Matt Edwards is the coordinator of musical theatre voice training and Artistic Director of the CCM Vocal Pedagogy Institute at Shenandoah University. He is the author of "So You Want to Sing Rock," a recipient of the Van Lawrence Fellowship, and has served as a master teacher for the NATS Intern program.

DISCLOSURES: None of the authors have conflicts of interest to disclose.

SINGING VOICE MACHINE LEARNING AI MODELS AS DIAGNOSTIC BIOMARKERS FOR VOCAL HEALTH, DYSPHONIA, TMD, AND TREMOR

THEODORA NESTOROVA^{1,2}, JIARUI XIE³, CHONGHUI ZHANG³, FELICIA NYANYO³, YAOYAO FIONA ZHAO^{3,4}, LUC MONGEAU⁴

¹ Schulich School of Music, McGill University, Montréal, Québec, Canada

² Department of Communication Disorders and Sciences, Viterbo University, La Crosse, Wisconsin, USA

³ Additive Design and Manufacturing Lab, McGill University, Montréal, Québec, Canada

⁴ Department of Biomedical and Mechanical Engineering, McGill University, Montréal, Québec, Canada

ABSTRACT

Introduction/Background: Vibrato—a multifactorial feature of singing voice function and expression—holds significant potential for assessing vocal health beyond its traditional aesthetic role. While often studied as an expressive device, recent research suggests vibrato may serve as a biomarker for neuromuscular dysfunction, biomechanical inefficiency, and motor control disorders. This study evaluates machine learning models for analyzing vibrato characteristics and their associations with biomechanical efficiency, muscle tension, and motor control abnormalities. By integrating advanced vibrato metrics, time-varying acoustic features, and physiological data, we aim to enhance diagnostic precision and bridge the gap between subjective assessment and objective, quantifiable voice analysis.

Methods: Multi-signal data were collected from 35 singers previously diagnosed with primary muscle tension dysphonia (pMTD) or temporomandibular disorder (TMD) through a collaborative voice care team at the McGill University Health Centre. Acoustic data originated from two prior vibrato studies recorded in identical environments to ensure reliability. Fifty expert judges—singing teachers and clinical voice specialists—rated vibrato samples for health, biomechanical efficiency, expressiveness, and regularity on a five-point Likert scale. Perceptual ratings were compared with acoustic measures, including Acoustic Voice Quality Index (AVQI), Cepstral Peak Prominence Smoothed (CPPS), and vibrato variability time-profiles (Nestorova, 2025). Physiological measures—motion capture (MOCAP) and surface electromyography (sEMG)—quantified jaw kinematics and muscle activation in the TMD subset. A convolutional neural network (CNN) baseline was optimized using Optuna-based hyperparameter tuning, varying convolutional layers (1–3), kernel size (2–4), hidden units (100–200), learning rate (10^{-5} – 10^{-1}), and dropout (0.01–0.1).

Results: Greater vibrato variability correlated significantly with increased muscle tension, particularly among singers with painful TMD. Confusion matrix analysis across training (N=600), validation (N=75), and test (N=76) datasets revealed distinct diagnostic separability: Healthy Controls achieved 83–90% classification accuracy, pMTD 67–69%, and painful TMD 23–30%. The CNN reached 100% accuracy on synthetic and 97–98% on real audio data, effectively differentiating healthy vibrato from dysfunction-associated irregularities, though performance varied by condition.

Conclusions: Vibrato-based AI analysis shows strong diagnostic promise. The relationship between stability, muscle tension, and neuromotor control highlights vibrato as a sensitive marker for pMTD, TMD, and emerging potential as a vocal tremor biomarker. While model precision was high, broader datasets and standardized protocols are needed for clinical and pedagogical translation. Refinement of feature extraction may establish vibrato as a key biomarker linking vocal artistry with neuromuscular health.

KEYWORDS: Vibrato; Machine Learning; Vocal Health; Diagnostic Biomarkers

PRESENTER BIO: Bulgarian-British-American voice scientist, pedagogue, and performer Theodora Ivanova Nestorova serves as Assistant Professor of Speech-Language Pathology at Viterbo University. A published author with international research awards and a former Fulbright Scholar, Theodora holds a PhD from McGill University, MBA from Global Leaders Institute, MM from New England Conservatory, and BM from Oberlin. More at theodoranestorova.com.

DISCLOSURES: None of the authors have conflicts of interest to disclose.

REFERENCES

Robin, J., Harrison, J. E., Kaufman, L. D., Rudzicz, F., Simpson, W., & Yancheva, M. (2020). Evaluation of speech-based digital biomarkers: Review and recommendations. *Digital Biomarkers*, 4(3), 99-108.

Bahr, R., Anibal, J., Bedrick, S., Bélisle-Pipon, J. C., Bensoussan, Y., Blaylock, N., ... & Rameau, A. (2024). Voice as an AI biomarker of health—introducing audiomics. *JAMA Otolaryngology–Head & Neck Surgery*, 150(4), 285-286. <https://doi.org/10.1001/jamaoto.2024.0068>

Fagherazzi, G., Fischer, A., Ismael, M., & Despotovic, V. (2021). Voice for health: The use of vocal biomarkers from research to clinical practice. *Digital Medicine*, 4(1), 78. <https://doi.org/10.1038/s41746-021-00447-2>

Nestorova, T. I. (2025). Vocal vibrato variability: Novel analytical tools & diagnostic-pedagogical approaches in diverse genres (Doctoral dissertation, McGill University).

Nestorova, T.I., Xie, J., Zhang, C., Nyanyo, F., Zhao, Y.F., Mongeau, L. (2025). Vibrato Machine Learning AI Models as Diagnostic Biomarkers for Vocal Health & Dysphonia. *NCVS Insights: AI Conference Version*. DOI forthcoming.

SIX-PACK SINGING: ON ASSOCIATIONS BETWEEN ANTERIOR ABDOMINAL MUSCLE ACTIVITY AND VITAL CAPACITY PERCENTAGE

CHRISTIAN T. HERBST^{1,2}, ANNA MILES³, JACQUELINE ALLEN^{3,4}, CALVIN P. BAKER^{3,5}

¹ Bioacoustics Laboratory, Department of Behavioral and Cognitive Biology, University of Vienna, Vienna, Austria

² Department of Communication Sciences and Disorders, College of Liberal Arts and Sciences, University of Iowa, Iowa City, IA, USA

³ Swallowing and Voice Research Laboratory, Speech Science, School of Psychology, University of Auckland, Auckland, New Zealand

⁴ Department of Surgery, Medical Health Sciences, University of Auckland, Auckland, New Zealand

⁵ School of Music, University of Auckland, Auckland, New Zealand

ABSTRACT

Breath management techniques are a key component within singing voice pedagogy, as well as clinical treatment of both hypo- and hyperfunctional voice disorders. Common assumptions regarding the relationships between anterior abdominal musculature (AAM), subglottal pressure (PSub), and sound pressure level (SPL) suggest a linear relationship where increased AAM activation will result in greater PSub causing a rise in SPL. The present study questions this tenet, utilizing simultaneous surface electromyographic, aerodynamic, and acoustic measures to contribute to this gap in understanding. Twenty-one healthy adult singers were recorded performing a slow vital capacity (VC) task and consonant-vowel repetitions (/papapa.../) from maximum inspiration to depletion, phonating on a controlled fundamental frequency at three SPLs. The results revealed a strong systematic increase in AAM activation as %VC decreased below 40% ($P < 0.0005$, $\rho = -0.81$). During controlled phonation, a strong %VC effect was also observed, with incrementally higher %sEMG values for respectively increasing SPL targets. Surprisingly, %sEMG was not strongly associated with either PSub or SPL, and immediate and sustained maximal contraction of the AAM throughout phonation did not significantly alter target vocal outputs compared with baseline measurements. The current study offers three key findings: 1) in habitual behaviors of healthy singers, AAM activates predominantly in association with decreasing %VC; 2) AAM activity is not strongly correlated with PSub or SPL; and 3) immediate contraction of AAM within the inspiratory reserve appears unnecessary for phonation. These data offer novel insights into relationships between respiratory and phonatory parameters and contribute to evidence-based voice pedagogy and therapy.

KEYWORDS: Breathing aerodynamics; Surface electromyography; Singing; Breath Support

PRESENTER BIO: Christian T. Herbst is an Austrian voice scientist and voice pedagogue. The focus of Christian's work is on basic voice science, singing voice physiology, and on mammalian voice production. He has to date published 80 peer-reviewed papers, including three publications in the prestigious Science journal. -- see <https://www.christian-herbst.org> for further details.

DISCLOSURES: None of the authors have conflicts of interest to disclose.

REFERENCES

Calvin Peter Baker, Anna Miles, Jacqueline Allen, Christian T. Herbst: "Six-Pack Singing: On Associations Between Anterior Abdominal Muscle Activity and Vital Capacity Percentage". In Press, Corrected Proof. DOI: 10.1016/j.jvoice.2025.03.039

KEYNOTE
**SOUNDING THE SELF: VOCALITY IN PERFORMANCE, PEDAGOGY,
AND CARE**

ALEXIS DAVIS-HAZELL^{1,2}

¹ Associate Professor of Voice, Assistant Director of Undergraduate Studies, University of Alabama School of Music,
Tuscaloosa, Alabama, USA

² President, National Association of Teachers of Singing

ABSTRACT

In an era increasingly attuned to critical inquiry of identity, embodiment, and cultural expression, the human voice endures as more than a biomechanical system; it is a dynamic site of meaning-making. Vocality is a framework that bridges physiology and identity, recognizing the voice as both an expressive instrument and a marker of social and cultural experience.

Drawing on interdisciplinary research, we will consider how vocal sound is shaped by intersecting factors such as self-perception, physical and cognitive capacities, demographic identifiers, and sociocultural contexts. These interwoven influences have profound implications for how voices are trained, perceived, and valued. We will consider how traditional vocal training paradigms can inadvertently or deliberately reinforce limiting normative ideals, and how a more expansive understanding of vocality can foster spaces for authenticity, agency, and inclusion.

This discussion invites singers, educators, clinicians, and researchers to foster a collaborative ecosystem which cultivates practices that honor the whole person behind the sound.

KEYWORDS: vocal self-perception, voice pedagogy, cultural expression

PRESENTER BIO: Alexis Davis-Hazell is an American mezzo-soprano known for her operatic and symphonic works in North America, over 130 performances of *Porgy and Bess* internationally, and contributing to the GRAMMY™-winning album *Grechaninov: Passion Week*. She is the National President of the National Association of Teachers of Singing and an Associate Professor at The University of Alabama.

DISCLOSURES: None of the authors have conflicts of interest to disclose.

THE VALSALVA EAR PRESSURE EQUALIZATION MANEUVER IN OPERATIC SINGERS: A QUALITATIVE PILOT STUDY

IAN DENOLFO, MM¹, GIOACCHINO LIVIGNI^{2,3}, DAVID MEYER, DM^{4,5,6},
ROBERT T. SATALOFF, MD, DMA, FACS^{7,8}

¹The Voice Foundation, Philadelphia, PA

²Curtis Institute of Music, Philadelphia, PA

³Royal Opera House (Covent Garden), Jette Parker Young Artists Programme, London, U.K.

⁴University of Iowa School of Music, The University of Iowa, Iowa City, Iowa, USA

⁵Communication Sciences and Disorders, The University of Iowa, Iowa City, Iowa, USA

⁶Shenandoah University, Winchester, Virginia, USA

⁷Drexel University College of Medicine, Dept of Otolaryngology – Head and Neck Surgery, Philadelphia, PA

⁸Sidney Kimmel Medical College Thomas Jefferson Univ, Dept of Otolaryngology, Philadelphia, PA

ABSTRACT

Objective: Anecdotal evidence suggests that some singers find benefit in using the Valsalva ear pressure equalization maneuver during operatic performances. The PI of the current study (IDN) observes a beneficial auditory change following this pressure equalization in his own singing, particularly in his high frequency auditory self-perception. The prevalence of this behavior in other singers, and its potential benefit have not been examined. This qualitative study investigates this practice in high level operatic singers.

Methods: This study was designed as an exploratory, qualitative, comparative study using semi-structured interviews with twenty operatic singers (Bunch/Chapman taxonomy levels 1-2) who self-report using a Valsalva ear pressure equalization maneuver while performing.

The following research questions will be explored:

- What are some perceived benefits of using the Valsalva maneuver while singing?
- In what settings did the singer feel the need to use this maneuver (practice room, stage, concert hall)?
- How frequently did the singer use the Valsalva?
- Were physical, acoustical or environmental factors involved in necessitating the use of the Valsalva maneuver (physical illness, acoustics of the space, environmental conditions such as humidity or allergens)?

Conclusion: The researchers hope to further our knowledge of opera singers' use of the Valsalva ear pressure equalization maneuver. The pedagogical implications of this maneuver and any perceived benefits will be discussed. Future studies may examine the functional basis of the Valsalva ear pressure equalization maneuver, and the many physiological and psychological challenges singers face in performing operatic repertoire.

KEYWORDS: Valsalva Maneuver; Acoustics; Singer's Formant Cluster

PRESENTER BIO: Ian DeNolfo, current Executive Director of The Voice Foundation, has appeared as a leading tenor in many of the world's top opera houses including The Metropolitan Opera, Teatro alla Scala, Deutsche Oper Berlin, The Washington Opera, and Los Angeles Opera. Ian is a graduate of The Juilliard School and The Curtis Institute of Music.

ACKNOWLEDGMENTS: Thanks for the immeasurable help and guidance from Dr. Robert T. Sataloff.

DISCLOSURES: None of the authors have conflicts of interest to disclose.

REFERENCES

- Gelfman DM. The Valsalva maneuver, set in stone. *Am J Med.* 2021;134(6):823-824.
- Bunch M, Chapman J. Taxonomy of singers used as subjects in scientific research. *J Voice.* 2000;14(3):363-369.
- Braun V, Clarke V. Using thematic analysis in psychology. *Qual Res Psychol.* 2006;3(2):77-101.
- Sundberg J. The science of the singing voice. DeKalb: Northern Illinois University Press; 1987.
- Kantarci T, Karatas E. Physiological aspects of the Valsalva maneuver. *Clin Otolaryngol.* 2016;41(5):535-542.

TOWARD EVIDENCE-BASED CHORAL CONDUCTING GESTURE

JEREMY N. MANTERNACH

School of Music/College of Education, University of Iowa, Iowa City, IA, USA

ABSTRACT

Teacher-conductors often have intense beliefs as to what constitutes a “good” conducting gesture. That said, various accomplished conductors’ firmly held beliefs are often at odds with one another. In many instances, conducting students are asked to unlearn much of what a previous instructor had taught them in favor of the “proper” technique. Holden (2003) pointed out that “any discussion of conducting technique can be problematic” (p. 3) because there are so many disparate opinions regarding its definition. So, is the adage true: the only thing two conductors can agree upon is the incompetence of a third?

This presentation will include an overview of research on choral conducting gesture and in the social and neurosciences to consider how particular gestures might affect singers’ voicing, either intentionally or unintentionally. The goal is not to suggest a collection of “correct” gestures. Rather, it will be to consider an evidence-based framework in which a conductor’s gestural language might support singers’ voicing and expressive goals.

KEYWORDS: Choral Conducting Gesture; Chorister Voicing; Imitation; Choral Pedagogy

PRESENTER BIO: Jeremy Manternach, Ph.D., is Associate Professor and Chair of Music Education at the University of Iowa, where he teaches undergraduate and graduate choral pedagogy, music education, and research courses. His research interests include choral conducting gesture and singer efficiency, choral and vocal acoustics, adolescent vocal development, and teacher voice use.

DISCLOSURES: None of the authors have conflicts of interest to disclose.

REFERENCES

- Holden, R. (2003). The technique of conducting. In J. A. Bowen (Ed.) *The Cambridge companion to conducting* (pp. 3-16). Cambridge University.
- Manternach, J. N. (2016). Effects of varied conductor prep movements on singer muscle engagement and voicing behaviors. *Psychology of Music*, 44, 574-586. doi: 10.1177/0305735615580357
- Manternach, J. N. (2012). The effect of conductor lip rounding and eyebrow lifting on singer lip and eyebrow postures: A motion capture study. *International Journal of Research in Choral Singing*, 4, 36-46.
- Manternach, J. N. (2012). The effect of varied conductor preparatory gestures on singer upper body movement. *Journal of Music Teacher Education*, 22(1), 20-34. doi: 10.1177/1057083711414428

TREMBLE CLEFS: THERAPEUTIC GROUP SINGING FOR INDIVIDUALS WITH PARKINSON'S DISEASE

DR. ABBEY L. DVORAK¹, DR. SUN JOO LEE²

¹ Music Therapy, The University of Iowa, Iowa City, Iowa, USA

² Music Therapy, University of Missouri-Kansas City, Kansas City, Missouri, USA

ABSTRACT

Parkinson's disease (PD) is a progressive condition that causes movement problems due to a loss of dopamine in the brain. Over time, PD can also affect speech, swallowing, mood, sleep, and memory. Some of the most frustrating symptoms of PD are voice and speech difficulties. PD often affects the muscles used for talking, making speech softer, slower, or harder to understand. While traditional treatments can help, they may be costly, have side effects, or lack the motivation and social support needed for adherence and long-term success. Music can be a powerful tool to help people with PD maintain and improve their voice because neural pathways for singing and speaking overlap. When people with PD sing, they practice controlling their breath, strengthening their voice, and improving their clarity—all within an engaging, motivating, and supportive group environment.

Tremble Clefs is a therapeutic singing group for individuals with Parkinson's disease, their caregivers, and family members. Weekly 90-minute sessions incorporate structured vocal exercises, sensorimotor synchronization, diaphragmatic breathing, and group singing to stimulate the motor speech system and enhance vocal function. Participants also benefit from the social connection and encouragement of peers facing similar challenges, all under the guidance of a trained music therapist who understands the unique needs of people with PD. This therapeutic environment fosters engagement and improves quality of life, regardless of musical background or experience.

The presenters synthesize research and share its application to Tremble Clefs. Participants will understand the program aims to improve health outcomes through singing, raise public awareness through community performances, and increase access to specialized care for individuals with PD. Tremble Clefs-Iowa will perform as part of this presentation to demonstrate the biopsychosocial benefits of therapeutic group singing. Professionals attending the conference will gain insight into how therapeutic singing can be integrated into voice care for neurodegenerative conditions, and how community-based programs can serve as powerful adjuncts to clinical treatment. Tremble Clefs exemplifies the transformative potential of singing when paired with empathy, creativity, and community engagement, and is a celebration of resilience, voice, and connection.

KEYWORDS: Music Therapy; Therapeutic Group Singing; Parkinson's Disease

PRESENTER BIOS: Dr. Abbey Dvorak is Associate Professor and Director of Music Therapy at the University of Iowa. She is a board-certified neurologic music therapist focusing on evidence-based music interventions to meet biopsychosocial needs of individuals with neurodegenerative conditions. Dr. Dvorak is current director of Tremble Clefs Iowa.

Dr. Lee is Assistant Professor of Music Therapy at the University of Missouri–Kansas City. A vocalist and board-certified neurologic music therapist, she focuses on therapeutic group singing for voice rehabilitation in Parkinson's disease. She has directed Tremble Clefs Arizona for 18 years and founded Tremble Clefs Iowa in 2023.

ACKNOWLEDGMENTS: We would like to thank the members of Tremble Clefs and their families, and the team members who collaborated on these projects. The authors would especially like to thank Sue and Frank Cannon and Kate Gfeller for their support of the music therapy program and our work in music and neurodegeneration.

DISCLOSURES: Funding provided by the Williams-Cannon Faculty Fellowship Award, the Iowa Mid-Career Faculty Scholar Award, the UI Office of Community Engagement, and individual and community donors.

REFERENCES

- Lee, S. J., Dvorak, A. L., & Manternach, J. (2024). Outcomes of singing and semi-occluded vocal tract exercises for individuals with Parkinson's disease. *Journal of Music Therapy*, 61(2), 132-167.
- Lee, S. J., & Dvorak, A. L. (2023). Therapeutic group singing for individuals with Parkinson's disease: A conceptual framework. *Music Therapy Perspectives*.

KEYNOTE
VOICE DISORDERS THAT PRESENT WITHOUT DYSPHONIA

THOMAS L. CARROLL, MD

Brigham and Women's Program for Voice, Swallowing and Upper Airway Health, Brigham and Women's Hospital,
Department of Otolaryngology-Head and Neck Surgery, Harvard Medical School, Boston, MA

ABSTRACT

This lecture will focus on the complaints that present to laryngologists and speech-language pathologists where dysphonia --- hoarseness, vocal fatigue, vocal effort --- is not the complaint; however, the problem ultimately is an issue with the larynx and most often glottic closure. It will help the attendee understand why some laryngeal pathologies can present with throat clearing, excessive mucus complaints, globus sensation and chronic cough. A state-of-the-art review of the modern workup and management of laryngopharyngeal reflux and glottic insufficiency diagnosis and management will be covered, including office based, permanent solutions for glottic insufficiency. This should help the listener understand why a normal flexible laryngoscopy exam is inadequate to rule in or rule out reflux or glottic insufficiency as the cause of the patient's complaint. The value of objective testing or empiric trials to treat reflux, other than acid suppression, and the use of diagnostic vocal fold augmentation and subsequent permanent, real-time stroboscopy guided office-based vocal fold augmentation will be presented.

KEYWORDS: Laryngeal Pathologies, Glottic Insufficiency, Laryngopharyngeal Reflux

PRESENTER BIO: Dr. Thomas Carroll directs the BWH Program for Voice, Swallowing and Upper Airway Health at Brigham and Women's Hospital, and is an Associate Professor of Otolaryngology-Head and Neck Surgery at Harvard Medical School. He is a laryngologist who focuses on professional voice care, office-based laryngeal surgery, reflux, and chronic cough.

DISCLOSURES: Dr. Carroll is a consultant for Pentax Medical, GSK, GH research, Value Analytics and Ambu. He serves on the scientific advisory board for and has stock options from Sofregen Medical and N-Zyme Biomedical. He receives royalties from Plural publishing.

DECODING ILLUSORY MELODY FROM ELECTROENCEPHALOGRAPHY USING BACKWARD TEMPORAL RESPONSE FUNCTIONS

AHYEON CHOI AND INYONG CHOI

Communication Sciences and Disorders, University of Iowa, Iowa City, IA, USA

ABSTRACT

Musical illusions provide a window into the constructive processes of auditory perception. The scale illusion, in which dichotically presented tones are reorganized into a coherent melody distinct from the physical input, exemplifies how the brain integrates and reconstructs sound. While previous neuroscience studies have examined such illusions using event-related potentials (ERPs) or spatial activation patterns, the temporally evolving neural representation of these perceptual reorganizations remains unclear. This study aimed to characterize the neural dynamics of the scale illusion in normal-hearing (NH) listeners using linear backward temporal response function (TRF) decoding of continuous EEG.

We recorded 64-channel EEG from 27 NH participants who listened to three variations of the scale illusion: (1) monaural presentation of the original dichotic stimulus, (2) monaural presentation of the perceived illusory pitch contour, and (3) the original scale illusion presented binaurally. TRF models were trained on the two monaural conditions (Stimuli 1 and 2) and tested on the binaural condition (Stimulus 3). After each trial, participants reported whether they perceived a continuous or fragmented melody. Pitch interval, amplitude envelope, and mel-frequency cepstral coefficient (MFCC) features were extracted and used as targets for EEG decoding.

Backward decoding revealed that, for the illusory (binaural) condition, MFCC achieved the highest correlation when trained on Stimulus 1 ($r = 0.193$), Pitch showed the highest decoding when trained on Stimulus 2 ($r = 0.098$), and Envelope performed relatively poorly ($r = 0.021$). These findings indicate that Pitch and MFCC features complementarily capture the neural tracking of illusory melodic organization, with Pitch-based decoding reflecting neural alignment to the perceived rather than the physical contour. In contrast, the Envelope feature alone was insufficient to account for neural reconstruction of the illusion, underscoring the importance of spectral and interval-based representations.

These results demonstrate that illusory melodies are not merely perceptual constructs but are robustly encoded in neural dynamics. Pitch and MFCC features jointly contribute to decoding the perceptual organization of the illusion, reflecting the brain's ability to reconstruct internally coherent melodic structures from ambiguous acoustic input. Furthermore, demonstrating that continuous stimuli can be encoded and decoded based on pitch features highlights the potential to extend this framework to explore neural encoding of singing voice stimuli in future research.

KEYWORDS: Musical illusion; Pitch; EEG; Neural Decoding; Temporal Response Function (TRFs)

PRESENTER BIO: Dr. Ahyeon Choi is a Postdoctoral Research Scholar working with Dr. Inyong Choi at the University of Iowa. Her research bridges music cognition and auditory neuroscience, focusing on EEG-based neural decoding of musical perception and illusions. She aims to elucidate how the brain encodes pitch, melody, and emotion in complex and naturalistic auditory contexts.

DISCLOSURES: None of the authors have conflicts of interest to disclose.

RETHINKING AUDITORY-PERCEPTUAL SCALING: ADVANCING RESEARCH AND CLINICAL IMPACT IN SPEECH AND VOICE ASSESSMENT

MILI KURUVILLA-DUGDALE

Department of Communication Sciences and Disorders, University of Iowa, Iowa City, Iowa, USA

ABSTRACT

Scaling methods are widely used to assess speech and voice function, despite well acknowledged limitations associated with their use, like poor listener reliability and task-related biases. The talk will highlight signal, rater, and task related challenges associated with conventional scaling approaches: the current methods are insufficient to address the multidimensionality of deviant speech and voice features that also fail to map onto known physical units. Higher order scales, when matched to appropriate psychophysical continua can provide a more reliable, valid, and interpretable framework for measurement. This talk advances the argument that psychophysical scaling principles offer a path forward for transforming perceptual speech and voice assessment.

The COVID-19 pandemic revealed both the challenges and opportunities for perceptual assessment. While resource constraints and research restrictions forced innovation in virtual auditory-perceptual studies, they also exposed how insufficient current practices are without standardized methods. Drawing from existing studies, the presentation will highlight what is and is not effective in improving perceptual scaling. The talk will cover recent work on scale fit, auditory training, and how the use of anchors all influence rater reliability and criterion and construct validity. The talk will differentiate the role of synthetic anchors, which allow precise control of isolated speech features, from natural anchors that preserve ecological validity. Both approaches can be leveraged to strengthen speech and voice assessment when paired with rigorous auditory training.

Reframing assessment through a psychophysical lens could result in a paradigm shift. It allows for the development of standardized and deployable tools that can scale across raters, timepoints, and settings. By optimizing perceptual scaling methods by aligning measurement level and psychophysical continua, providing anchors, and implementing auditory training, we can enhance scientific rigor and clinical relevance. This work sets the stage for a new direction in speech and voice assessment characterized by a standardized approach, where perceptual science and clinical translation converge to improve research and clinical outcomes.

KEYWORDS: e.g. Dysarthria, Measurement level, Psychometric properties, Scales

PRESENTER BIO: Mili Kuruvilla-Dugdale is an Associate Professor at the University of Iowa. Her research focuses on improving foundational knowledge about speech physiology in neurodegenerative diseases, with the goal of developing more sensitive assessments. Another line is focused on developing optimal perceptual scaling methods that are grounded in psychophysical principles and have strong psychometric properties.

ACKNOWLEDGMENTS: Funding ACKNOWLEDGMENTS: American Speech-Language-Hearing Foundation, New Century Scholars Grant (PI: Kuruvilla-Dugdale) and National Institute on Deafness and Other Communication Disorders, R21DC019952 (mPI: Kuruvilla-Dugdale).

DISCLOSURES: Financial disclosure: Full-time tenure-track faculty member at the University of Iowa.

AI-GENERATED DYSPHONIC SPEECH: MORE CLEAR, LESS REAL?

**PASQUALE BOTTALICO¹, CHARLES J. NUDELMAN^{1,2}, DANIEL FOGERTY¹, VIRGINIA TARDINI^{1,3},
KEIKO ISHIKAWA⁴**

¹ Department of Speech and Hearing Science, University of Illinois Urbana-Champaign, Champaign, IL, US

² Department of Communication Sciences & Disorders, Syracuse University, Syracuse, NY, US

³ Department of Industrial Engineering, University of Bologna, Bologna, Italy

⁴ Department of Communication Sciences & Disorders, University of Kentucky, Lexington, KY, US

ABSTRACT

Artificial Intelligence (AI)-driven voice cloning has rapidly advanced, offering new possibilities for augmenting research on voice disorders where clinical data are scarce. Dysphonia, a condition marked by impaired voice quality, often leads to reduced speech intelligibility in noise and diminished quality of life. However, studying its impact is challenging due to limited datasets and variability across pathologies. This work investigates whether AI-generated dysphonic voice clones can accurately reproduce the perceptual characteristics of disordered speech, particularly regarding intelligibility.

We conducted three experiments using recordings from 12 speakers (six dysphonic, six vocally healthy) and their AI-generated voice clones. Sixty-four listeners evaluated discrimination, identification, and intelligibility of speech samples. In the first experiment, listeners judged whether paired samples were from the same speaker. Sensitivity was higher for dysphonic than normal voices, indicating that synthetic dysphonic voices exhibited perceptual differences more detectable to listeners. The second experiment tested identification accuracy of AI-generated versus natural speech. Performance was modest overall, with higher accuracy when comparing real-to-real voice pairs and improved detection for dysphonic voices compared to normal voices.

The third experiment assessed intelligibility in noise using the Hearing-in-Noise Test. Strikingly, AI-generated dysphonic voices were consistently more intelligible than natural dysphonic voices. For male speakers, intelligibility improved from 35.5% (real) to 66.5% (AI), while female dysphonic voices reached 67.0% intelligibility in the AI-generated condition. Acoustic analyses revealed that cloned voices reduced spectral variability and enhanced clarity, features associated with improved intelligibility but inconsistent with the degraded signal of true dysphonia.

These findings highlight both opportunities and limitations of current AI-based cloning systems. While AI-generated dysphonic voices may serve as useful data-augmentation tools and even support assistive communication applications, they fail to replicate the reduced intelligibility characteristic of natural dysphonic speech. This limitation underscores the need for improved synthesis models capable of accurately capturing pathological voice quality.

Overall, the study demonstrates the potential of generative AI in clinical speech research while emphasizing caution in its application. The results have implications for speech intelligibility modeling, clinical training, and the development of more representative datasets for studying voice disorders.

KEYWORDS: e.g. Voice cloning; Dysphonia; Speech intelligibility; Artificial Intelligence

PRESENTER BIO: Pasquale Bottalico, an Associate Professor in the department of Speech and Hearing Sciences at the University of Illinois-UC. With degrees in telecommunications engineering and opera singing, his work bridges science and performance. His research focuses on classroom acoustics, vocal load, and speech intelligibility. He holds a Ph.D. in Metrology and performs under renowned conductors worldwide.

ACKNOWLEDGMENTS: The authors gratefully acknowledge Gonzalo Farid Saud Medina, Kate Harty, Kaliyah House-Henry, Ryan Anderson, Naomi Ha, Emma Morrical, and Aron Olivares for their assistance and contribution to this research study.

DISCLOSURES: None of the authors have conflicts of interest to disclose.

ASSESSMENT OF VOCALIZATION IN A RAT AND FINDINGS FROM THE *PINK1*^{-/-} RAT MODEL OF PARKINSON DISEASE

MARYANN N. KRASKO, PhD^{1,2,3,4}, DENIS MICHAEL RUDISCH, MM, DIPL UNIV^{1,2,5}, MICHELLE R. CIUCCI, PhD, CCC-SLP^{1,2,6}

¹ Department of Otolaryngology-Head and Neck Surgery, University of Wisconsin School of Medicine and Public Health, Madison, WI, USA

² Department of Communication Sciences and Disorders, University of Wisconsin-Madison, Madison, WI, USA

³ Department of Otolaryngology, University of Iowa, Iowa City, IA, USA

⁴ Department of Communication Sciences and Disorders, University of Iowa, Iowa City, IA, USA

⁵ Institute for Clinical and Translational Research, University of Wisconsin-Madison, Madison, WI, USA

⁶ Neuroscience Training Program, University of Wisconsin-Madison, Madison, WI, USA

ABSTRACT

Vocal communication is a critical but largely understudied area in the field of neurological disease. In Parkinson disease (PD), vocal deficits are among the earliest and most prevalent signs,^{1,2} often emerging in the prodromal stage well before the onset of hallmark motor signs.² Despite this, these vocal changes remain underrecognized and undertreated in the clinic, in part due to limited understanding of their early manifestations and sex-specific trajectories. Rat ultrasonic vocalizations (USVs) offer a translationally relevant tool for modeling and assessing these early vocal changes,^{3,4} as they share key vocal acoustic features with human vocalization. This talk will begin with a brief overview of methodologies used to study USVs in rats, highlighting their utility as behavioral biomarkers with potential for early detection and therapeutic targeting.

To illustrate this application, I will present data from the *Pink1* knockout (*Pink1*^{-/-}) rat, a validated genetic model of early-onset PD. Sixty-four rats (n=16 per genotype/sex group) were tested at 4 months of age, a timepoint analogous to the prodromal stage of disease. Calls were elicited via mating paradigm, after which USVs were recorded for 90 seconds. DeepSqueak software was used to categorize calls and to extract parameters including call count (#), call length (ms), average principal, high, low, and peak frequencies (kHz), bandwidth (kHz), and intensity (dB/Hz). Findings revealed multiple sex-specific genotype effects. Male *Pink1*^{-/-} rats exhibited significantly higher low frequencies in both simple ($p=0.034$) and frequency-modulated (FM) ($p=0.0002$) calls compared to male wild-type (WT) controls. In contrast, female *Pink1*^{-/-} rats showed significantly lower peak ($p=0.032$) and low ($p=0.034$) frequencies compared to female controls. Interestingly, while WT males and females displayed significant differences across nearly all vocal acoustic measures ($p<0.05$), these sex-based distinctions were largely absent in the *Pink1*^{-/-} rats. Call count, length, and bandwidth did not differ significantly by genotype or sex ($p>0.05$).

These findings suggest that early deviations in vocal acoustics may occur in a sex-specific manner, with males and females demonstrating distinct frequency-related changes in the prodromal stage of PD. Moreover, the apparent loss of expected sex differences in *Pink1*^{-/-} rats may reflect early pathological disruption of normal vocal behavior. Together, this work supports the use of USVs as sensitive translational markers of early PD-related dysfunction and underscores the importance of considering sex as a biological variable. By advancing our understanding of prodromal vocal changes, this research contributes to the development of behavioral biomarkers with the potential to improve early diagnosis and intervention strategies in PD.

KEYWORDS: Voice; Vocalization; Parkinson disease; Modeling

PRESENTER BIO: Dr. Maryann Krasko is a postdoctoral fellow in the Department of Otolaryngology at the University of Iowa Health Care. She received a PhD from the University of Wisconsin-Madison, an MS from Columbia University, and a BS from New York University. Her primary areas of interest are concerned with the mechanisms underlying upper airway function and disorders.

ACKNOWLEDGMENTS: This work was funded by the National Institute on Deafness and Other Communication Disorders (F31DC022161; T32DC009401; R01DC018584).

DISCLOSURES: None of the authors have conflicts of interest to disclose.

REFERENCES

Ho, A.K.; Iansek, R.; Marigliani, C.; Bradshaw, J.L.; Gates, S. Speech impairment in a large sample of patients with Parkinson's disease. *Behav. Neurol.* 1998, 11, 131–137, doi:10.1155/1999/327643.

Ma, A., Lau, K. K., & Thyagarajan, D. (2020). Voice changes in Parkinson's disease: What are they telling us?. *Journal of Clinical Neuroscience*, 72, 1-7.

Grant, L. M., Kelm - Nelson, C. A., Hilby, B. L., Blue, K. V., Paul Rajamanickam, E. S., Pultorak, J. D., ... & Ciucci, M. R. (2015). Evidence for early and progressive ultrasonic vocalization and oromotor deficits in a PINK1 gene knockout rat model of Parkinson's disease. *Journal of neuroscience research*, 93(11), 1713-1727.

Krasko, M. N., Hoffmeister, J. D., Schaen-Heacock, N. E., Welsch, J. M., Kelm-Nelson, C. A., & Ciucci, M. R. (2021). Rat models of vocal deficits in Parkinson's disease. *Brain Sciences*, 11(7), 925.

UNISONO AFTER LARYNGECTOMY: FINDING A NEW VOICE THROUGH MUSIC AND SCIENCE

THOMAS MOORS^{1,2}, DAVID HOWARD³, EVANGELOS HIMONIDES⁴, DAVID MEYER^{5,6,7}

¹ Director, Shout at Cancer; ² ENT Department, Lancashire Teaching Hospitals, Lancashire, United Kingdom

³ Professor Emeritus, Dept of Electronic Engineering, Royal Holloway, Univ of London, United Kingdom

⁴ Professor of Technology, Education, and Music, University College London, London, United Kingdom

⁵ School of Music, University of Iowa, Iowa City, Iowa, USA

⁶ Department of Communication Sciences and Disorders, University of Iowa, Iowa City, Iowa, USA

⁷ Shenandoah Conservatory, Winchester, Virginia, USA

ABSTRACT

Laryngectomy—the surgical removal of the larynx—is a life-saving procedure, but one that profoundly alters a person’s life. Beyond the loss of natural voice, patients face a cascade of physical and psychosocial challenges that affect communication, identity, social interaction, and emotional well-being. Many individuals are left isolated, struggling to reintegrate into society.

In this context, music has emerged as an unexpected yet powerful tool for rehabilitation. Through singing, breathing, beatboxing, rhythm, and acting exercises, individuals who have undergone laryngectomy are rediscovering not only how to speak, but how to express themselves. The Shout at Cancer Laryngectomy Choir offers a unique model of recovery that bridges clinical rehabilitation with creative practice. Working alongside voice coaches, scientists, and artists, participants engage in co-creative workshops that blend music, poetry, and science to strengthen respiratory control, articulation, timing, and confidence.

Beyond functional gains, the psychosocial impact of these musical interventions is profound. The collective act of singing fosters belonging, mutual understanding, and emotional resilience. Participants report improved mood, increased social connectedness, and renewed motivation. The choir’s performances go even further: they challenge public perceptions of life after laryngectomy, offering audiences direct insight into the realities and possibilities of communication without a voice box.

Each performance becomes a platform for advocacy and education, engaging not only the public but also researchers, clinicians, engineers, and policymakers. The choir thus serves as both a therapeutic community and a catalyst for interdisciplinary collaboration, inspiring advances in speech technology, surgical care, and post-operative support.

Through this initiative, individuals who once faced silence are reclaiming their identity as active contributors to science, art, and society. Their voices now carry a message that transcends medicine: that creativity can reframe disability as resilience, and rehabilitation as art.

This shared voice finds its most powerful symbol in the Laryngectomy Vocal Tract Organ—an instrument constructed from 3D-printed MRI scans of the singers’ reconstructed vocal tracts. In this organ, every interdisciplinary voice that contributed to the project—patients, clinicians, scientists, engineers, and artists—literally comes together and resonates as one. It embodies the project’s essence: that healing is not only the restoration of function, but the harmonisation of human creativity, technology, and compassion into a new, collective sound.

KEYWORDS: Laryngectomy; Voice restoration; Music-based rehabilitation; Psychosocial adaptation; Patient engagement;

PRESENTER BIO: Thomas Moors is a Belgian, UK-based award-winning medical doctor, PhD student at Ghent University, and founder of Shout at Cancer. He integrates art into healthcare, pioneering music-based speech rehabilitation after laryngectomy. A grant expert and voice coach, he also serves as Medical Director of Sound Voice and board member of The Listening Planet Foundation.

DISCLOSURES: None of the authors have conflicts of interest to disclose.

KEYNOTE
THE BIOMECHANICS OF VOCAL FOLD TRAUMA

JACK J. JIANG, MD, PhD

Department of Surgery, Division of Otolaryngology
Laryngeal Physiology Lab
University of Wisconsin-Madison, Madison, Wisconsin, USA

ABSTRACT

This presentation explores the mechanisms of vocal fold injury and structural alteration through computational modeling. It begins with an introduction to the two-mass model, providing both mathematical and physical justification for its effectiveness in representing the vocal folds as a coupled spring-mass system. The discussion then examines how collision and contact forces contribute to phonotraumatic damage, integrating insights from both modeling outcomes and experimental data. Next, the presentation addresses fluid redistribution within the vocal fold and its implications for tissue deformation and lesion formation. An overview of finite element methods (FEM) follows, leading to computational results that illustrate the fluid-structure interactions involved in nodule and polyp etiology. Structural features of the vocal fold along with treatment options for nodules and lesions are then briefly discussed: followed by conclusions drawn from modeling. Collectively, these modeling approaches enhance our understanding of how repetitive mechanical stress and fluid dynamics drive the onset of vocal fold pathologies.

KEYWORDS: Vocal fold biomechanics; Phonotrauma; Vocal fold injury

PRESENTER BIO:

Dr. Jack J. Jiang, MD, PhD, is a Professor at the University of Wisconsin-Madison specializing in otolaryngology. He co-directs the Voice Research Training Program and leads the Otolaryngic Biomedical Engineering Research Center. His research focuses on vocal fold biomechanics, laryngeal physiology, and voice disorder assessments.

DISCLOSURES: None of the authors have conflicts of interest to disclose.

EARLY DETECTION OF VOICE SYMPTOMS IN PRE-SERVICE TEACHERS OF SOUTH AMERICA

LADY CATHERINE CANTOR-CUTIVA¹, LÉSLIE PICCOLOTTO FERREIRA², MARIA LÚCIA VAZ MASSON³, IARA BITTANTE DE OLIVEIRA⁴, MARIA CELINA MALEBRAN BEZERRA DE MELLO⁵, ALEJANDRO MORALES QUEVEDO⁶, CARLOS MANZANO⁷, MARÍA DEL CARMEN DALMASSO⁸, ADRIANA DÍAZ GUTIÉRREZ⁹, JESSICA LAURA RAMONDA¹⁰, ERIC J. HUNTER¹¹

¹Department of Audiology & Speech-Language Pathology, East Tennessee State University, Johnson City, TN, USA

²Pontifícia Universidade Católica de São Paulo, São Paulo, Brasil

³Universidade Federal da Bahia, Bahia, Brasil

⁴Pontifícia Universidade Católica de Campinas, Campinas, Brasil

⁵BonMed, Santiago de Chile, Chile

⁶Universidad Manuela Beltrán, Bogotá D.C., Colombia

⁷Hospital Médica Sur/Centro Médico ABC, Ciudad de Mexico, Mexico

⁸Universidad Nacional del Oeste, Buenos Aires, Argentina

⁹Dirección Nacional de Sanidad Policial / Hospital Policial, Montevideo, Uruguay

¹⁰Instituto superior ISOF / Hospital Privado de Córdoba, Córdoba, Argentina

¹¹Department of Communication Sciences and Disorders, University of Iowa, Iowa City, IA, USA

ABSTRACT

Objective: To determine the prevalence and associated factors of vocal symptoms in pre-service teachers.

Methods: A multicenter cross-sectional study involving undergraduate students from Argentina, Brazil, Chile, Colombia, Mexico, and Uruguay. Data collection on training conditions and vocal symptoms of pre-service teachers was conducted using the Prospective Teacher Voice Questionnaire (PTVQ). Dependent variables included the factors of the Vocal Fatigue Index (VFI), the total score of the Voice Symptom Scale (VoiSS), and the Voice Disorder Screening Index (SIVD). Independent variables were country of origin, year of training, current participation in academic practice, and gender.

Results: A total of 1,666 undergraduate students participated in the study (15 from Brazil, 22 from Argentina, 98 from Chile, 145 from Colombia, 165 from Uruguay, and 1,221 from Mexico). Country-specific analysis showed that 82% of Argentine students, 73% of Brazilian, 62% of Chilean, 59% of Colombian, 50% of Uruguayan, and 31% of Mexican students reported being in academic practice. Regarding dependent variables, the prevalence of vocal fatigue (Factor 1 of the VFI) was 43% (n=714), the prevalence of vocal symptoms (VoiSS) was 26% (n=425), and the prevalence of voice disorders (SIVD) was 14% (n=225). Analysis of the prevalence of vocal symptoms associated with academic practice suggests that 34% of students in practice reported vocal symptoms (OR=1.89, p-value 0.00). Country-specific analysis showed that Brazilian students had significantly more voice symptoms identified by the VoiSS (67% Brazil, 34% Chile, 32% Colombia, 27% Uruguay and Argentina, and 24% Mexico). Interestingly, although students in academic practice had a higher prevalence of voice symptoms in all countries, these differences were statistically significant only in Chile, Colombia, and Mexico.

Conclusion: The results suggest a significant effect of practice exposure and contextual factors on vocal symptoms in pre-service teachers. Argentina, with the highest participation in practice (82%) and a low prevalence of symptoms (27%), stands out for offering workshops on voice use during teacher training. In contrast, Brazil shows high participation in practice (73%) and the highest prevalence (67%). These differences suggest that vocal demands depend not only on practice time but also on cultural and environmental factors such as classroom noise, student behavior, and institutional support. Future research should aim to balance sample sizes across countries and study the impact of institutional initiatives on vocal care in teacher training. It is important to note that the unequal sample sizes (ranging from 22 to 1,221 participants) limit generalizability and may have introduced comparative bias.

KEYWORDS: pre-service teachers; voice symptoms; occupational voice

PRESENTER BIO: Assistant Professor in the Department of Audiology & Speech-Language Pathology at East Tennessee State University, Dr. Cantor Cutiva is the Director of the VoCES Research Lab. She holds a Master's degree in Health and Safety at Work, a second Master's degree, and a Ph.D. in Health Sciences.

ACKNOWLEDGMENTS: We would like to thank the many participants. The analysis protocols and techniques reported in this publication were partially supported by the National Institute of Deafness and Other Communication Disorders of the National Institutes of Health under Award Number R01DC012315 (P.I. Eric Hunter). The content is solely the responsibility of the authors and does not necessarily represent the official views of the National Institutes of Health.

DISCLOSURES: None of the authors have conflicts of interest to disclose.

REFERENCES

- Cantor-Cutiva, L. C., Malebran, M. C., Morales, A., Diaz, A., Cerda, F., & Hunter, E. J. (2025). Voice symptoms among South American prospective teachers. Associated Factors and Direct Costs. *Folia phoniatrica et logopaedica*.
- Cantor-Cutiva, L. C., Manzano, C., & Hunter, E. J. (2025). Prevalence and Associated Factors of Voice Symptoms among Mexican prospective teachers: A cross-sectional study. *Journal of Communication Disorders*, 106531.

MULTI-MODAL MRI AND CT IMAGE PROCESSING FOR VOCAL TRACT AREA FUNCTION ESTIMATION

SWATI RAMTILAK¹, DAVID MEYER^{2,5,6}, AIMING LU⁴, JAMES HOLMES^{3,1}, JARRON ATHA³,
ERIC HOFFMAN^{3,1}, SAJAN GOUD LINGALA^{1,3}

¹ Roy J. Carver Department of Biomedical Engineering, University of Iowa, Iowa City, Iowa

² School of Music, University of Iowa, Iowa City, Iowa

³ Department of Radiology, University of Iowa, Iowa City, Iowa

⁴ Medical Physics Division, Department of Radiology, Mayo Clinic

⁵ Communication Sciences and Disorders, University of Iowa, Iowa City, Iowa

⁶ Shenandoah University

ABSTRACT

Magnetic Resonance Imaging (MRI) based vocal tract models offer a comprehensive way of studying vowel production, swallowing movements and study voice quality utilizing the tract shape and area functions of the volume. Teeth form an integral part of these models altering the vocal tract shape near the lips. Conventional gradient echo MRI (GRE) sequences are widely used to generate vocal tract models which are not inclusive of teeth. Zero echo time (ZTE) MRI has demonstrated utility in imaging bony structures. This study aims at comparing vocal tract models generated using solely GRE-MRI with two hybrid models i.e. ZTE-MRI and CT-MRI that are teeth inclusive.

Teeth and mandible derived from CT scans are thresholded and manually segmented using 3D Slicer. ZTE scans are post processed² to highlight bone structure of the upper airway pertinent to speech production. The algorithm is a histogram log-based bias correction and a normalization scheme which was implemented on the raw ZTE images. The teeth and mandible are segmented manually using 3D Slicer and up sampled to match the CT volume and spacing. The bone structures in the upper airway from both CT and ZTE are registered using a rigid body landmark scheme to the GRE volume. The vocal tract soft tissue is segmented from the up sampled GRE scan and combined with ZTE and CT derived bone structures to generate hybrid vocal tract models.

Vocal tract area functions quantify changes in airway cross-sectional area along the tract from the glottis to the lips. Three professional voice users performed the speech task of uttering the vowel /a/ (as in father), and area functions were derived from the three vocal tract models. Using a fixed airway centerline (L0 at the glottis to L1 at the lips), 400 oblique midsagittal slices were taken, and the airway area in each slice was computed. Plotting cross-sectional area versus distance from the glottis yields a characteristic curve of vocal tract shape during /a/ vowel production.

GRE-MRI based vocal tract models tends to overestimate cross sectional area of the tract behind the lips as they omit the inclusion of teeth when compared to the hybrid ZTE or CT models. Thus ZTE-MRI and CT-MRI based hybrid models offer a comprehensive way to model vocal tracts from patient derived scans. ZTE-MRI offers an innate advantage over CT in not exposing patients to radiation.

KEYWORDS: Vocal Tract Models; Teeth; ZTE-MRI

PRESENTER BIO: Swati Ramtilak, B.S., M.S., is a PhD candidate in Biomedical Engineering at the University of Iowa, conducting research in Dr. Sajan Goud Lingala's lab.

ACKNOWLEDGMENTS: This work was supported by NIH under grant NIH NHLBI: R01 HL173483 and University of Iowa Office of Vice President's jump start P3 award. This work was conducted on an MRI instrument funded by NIH-S10 instrumentation grant: 1S10OD025025-01.

DISCLOSURES: None of the authors have conflicts of interest to disclose.

REFERENCES

Story, Brad H. "A parametric model of the vocal tract area function for vowel and consonant simulation." *The Journal of the Acoustical Society of America* 117.5 (2005): 3231-3254.

Wiesinger, Florian, et al. "Zero TE MR bone imaging in the head." *Magnetic resonance in medicine* 75.1 (2016): 107-114.

A 3D UPPER AIRWAY ATLAS: DATA, SEGMENTATIONS, AND PRINTABLE VOCAL TRACTS

SUBIN ERATTAKULANGARA^{1,2}, SWATI RAMTILAK¹, KARTHIKA KELAT¹, DAVID MEYER^{3,4,5},
SAJAN GOUD LINGALA^{1,6}

¹ Roy J. Carver Department of Biomedical Engineering, University of Iowa, Iowa City, Iowa, USA

²Department of medical physics, Memorial Sloan Kettering Cancer Center, New York, NY, USA

³Janette Ogg Voice Research Center, Shenandoah University, Winchester, Virginia, USA

⁴School of Music, University of Iowa, Iowa City, Iowa, USA

⁵Department of Communication Sciences and Disorders, University of Iowa, Iowa City, Iowa, USA

⁶Department of Radiology, University of Iowa, Iowa City, Iowa, USA

ABSTRACT

The human upper airway is vital for physiological functions like respiration and phonation, making its study crucial for diagnosing conditions such as obstructive sleep apnea [1], aiding surgical planning, and informing speech therapy [2]. Magnetic Resonance Imaging (MRI) is the preferred, non-invasive method for this study [3], providing structural and dynamic soft-tissue detail. The utility of 3D static MRI in generating Finite Element Method (FEM) models is clear [4], yet its application is hampered by the need for manual segmentation, a process that is notoriously time consuming and costly, requiring significant expert hours per data volume. While deep learning offers a potential solution for segmentation, its effectiveness is limited by the scarcity of large, labeled, domain-specific datasets, despite advances like transfer learning. This critical gap necessitates expert annotated data, yet the lack of robust automated tools means creating these datasets remains a labor intensive endeavor. Our proposed work directly addresses these challenges by making two significant contributions to the speech MRI community. Firstly, it provides a unique, publicly available dataset comprising volumetric MRI data from 20 English speakers, including six static voiced images per subject, their expert segmentations, and 3D printable vocal tracts. This dataset is foundational for developing and validating new state-of-the-art segmentation algorithms. Secondly, it introduces a no-code segmentation tool built using Slicer, Monal, Label, and Docker, designed to enable non-technical experts to generate preliminary segmentations, thus streamlining the dataset creation process. By making both the code and model weights open source, this research offers accessible resources that will significantly accelerate the advancement of airway analysis and facilitate the creation of extensive datasets in the field.

KEYWORDS: Voice; Deep Learning; Segmentation

PRESENTER BIO: Dr. Subin Erattakulangara is an expert in biomedical engineering, deep learning, and advanced MRI analysis. His research focuses on developing AI tools for accurate image segmentation of the vocal tract and rapid and free-breathing MRI reconstruction. His work enhances diagnostic precision for conditions like sleep apnea and improves patient comfort.

ACKNOWLEDGMENTS: We'd like to acknowledge Dr. David Meyer for providing the extensive manual segmentations necessary for a large number of volumetric tasks. His work served as the cornerstone for developing our automated AI tool.

DISCLOSURES: None of the authors have conflicts of interest to disclose.

REFERENCES

Daley, J. T. & Schwab, R. J. Imaging the airway in obstructive sleep apnea. in *Neuroimaging of Sleep and Sleep Disorders* (eds. Nofzinger, E., Thorpy, M. J. & Maquet, P.) 256–263 (Cambridge University Press, Cambridge, 2013). doi:DOI: 10.1017/CBO9781139088268.035.

Wong, A.-M. *et al.* Examining the impact of multilevel upper airway surgery on the obstructive sleep apnoea endotypes and their utility in predicting surgical outcomes. *Respirology* **27**, 890–899 (2022).

Darquenne, C. *et al.* Upper airway dynamic imaging during tidal breathing in awake and asleep subjects with obstructive sleep apnea and healthy controls. *Physiol Rep* **6**, e13711 (2018).

Xu, C., Brennick, M. J., Dougherty, L. & Wootton, D. M. Modeling upper airway collapse by a finite element model with regional tissue properties. *Med Eng Phys* **31**, 1343–8 (2009).

TRANSFORMING VOCAL VARIABILITY FROM NOISE TO SIGNAL: A PRECISION MEDICINE APPROACH TO VOICE HEALTH

ERIC J. HUNTER, MARK L. BERARDI

Department of Communication Science and Disorders, University of Iowa, Iowa City, Iowa, USA

ABSTRACT

Voice professionals, clinicians, and researchers have long struggled with variability in vocal measurements, often dismissing fluctuations as “noise” that obscure meaningful assessment. This view overlooks a fundamental insight: variability itself may carry diagnostic information about vocal function and health, much as variability analysis in cardiology and orthopedics has revealed fluctuation patterns that serve as sensitive indicators of resilience and emerging dysfunction. Yet early acoustic metrics such as jitter and shimmer, while initially promising, proved unstable across contexts and failed to capture broader patterns that may hold diagnostic value.

Vocal variability arises from four interacting domains: measurement systems, methodological protocols, physiological processes, and neurocognitive factors. Recent studies illustrate how these domains intersect to shape measurable voice behavior across multiple timescales and contexts. Psychosocial stress, for instance, can alter pitch and loudness variability alongside physiological changes such as heart rate and skin conductance. Similarly, variations in room acoustics, background noise, hydration, hormonal state, or fatigue elicit systematic adjustments in vocal effort and spectral balance. The central difficulty lies in disentangling variability arising from measurement or methodological artifacts from that reflecting genuine psychophysiological responses, and in identifying which of these patterns carry functional or physiological meaning. Although such fluctuations may appear stochastic, many represent lawful, adaptive adjustments through which the voice system balances stability and flexibility under changing internal and external demands. The challenge and opportunity is to separate meaningful variability from true uncertainty, recognizing that both can yield insight into system resilience and vulnerability.

To guide such interpretation, we propose a precision voice framework organized around four constructs: capacity (available physiological and psychological resources), demand response (adaptive strategies under load), reserve (the margin between current demand and maximum capacity), and recovery (restoration time courses). Advancing precision voice care through this framework will require repeated sampling to establish person-specific baselines, transparent methodological standards, and scalable infrastructures that support longitudinal and multi-site research. With this approach, voice science can move beyond snapshot-based assessment toward dynamic, systems-level evaluation that treats both structured adaptation and inherent uncertainty as integral to vocal function and health.

KEYWORDS: precision medicine; vocal variability; individualized assessment

PRESENTER BIO: Eric Hunter serves as Chairperson of Communication Sciences and Disorders at the University of Iowa and holds the Harriet B. and Harold S. Brady Chair. His NIH-supported research focuses on occupational voice use, particularly voice disorders in teachers, biomechanical modeling of vocal systems, and signal processing. He is a fellow of the Acoustical Society of America.

ACKNOWLEDGMENTS: This research was supported by the NIDCD of the NIH under Grant No. R01DC012315. The content is solely the responsibility of the authors and does not necessarily represent the official views of the National Institutes of Health.

DISCLOSURES: None of the authors have conflicts of interest to disclose.

REFERENCES

Hunter, E. J., & Berardi, M. L. (In Press). Embracing variability: Toward proactive and precision-based voice science. *Logopedics Phoniatrics Vocology*. <https://doi.org/10.1080/14015439.2025.2562929>

FROM SIGNAL VARIABILITY TO VOCAL SUSTAINABILITY: OPERATIONALIZING PRECISION MEDICINE FOR OCCUPATIONAL VOICE HEALTH

MARK L. BERARDI, ERIC J. HUNTER

Dept Communication Science and Disorders, University of Iowa, Iowa City, Iowa, USA

ABSTRACT

Vocal variability has long been dismissed as measurement noise, but emerging evidence demonstrates that structured fluctuations reflect adaptive capacity and individual differences in vocal function. Building on this reconceptualization of variability as diagnostic signal, we extend the four pillars of precision medicine—Predictive, Preventive, Personalized, and Participatory (P4)—to occupational voice health, demonstrating how individual variability patterns can inform proactive, individualized care.

Traditional voice assessment relies on single-measure snapshots that miss critical information about how individuals adapt to vocal demands. A precision approach instead focuses on functional voice evaluation: measuring individual capacity across multiple dimensions, quantifying real-time responses to vocal challenges, and tracking how the margin between capacity and demand changes during and after vocal loading. This shift enables clinicians to identify unsustainable vocal workloads before symptoms develop, moving from reactive treatment to proactive workload management.

Predictive capabilities emerge from recognizing that individuals show consistent, measurable patterns in how they respond to vocal challenges. Recent evidence demonstrates distinct compensation strategies—some speakers actively increase effort to maintain voice quality under load, while others economize effort by accepting quality changes. These individual demand response signatures enable risk stratification and early identification of workers at elevated risk for voice disorders.

Preventive interventions can be targeted based on these individual patterns rather than waiting for fatigue or pathology to manifest. Individuals showing signs of unsustainable workload patterns—where vocal demands consistently approach or exceed their adaptive capacity—can receive intensive capacity-building programs and environmental modifications, while those maintaining adequate safety margins receive standard wellness protocols.

Personalized assessment and treatment require establishing person-specific baselines rather than relying solely on normative data. Interventions can then be matched to individual compensation strategies, vocal demands, and recovery patterns, recognizing that what constitutes sustainable workload varies substantially across individuals.

Participatory engagement involves occupational voice users as active collaborators in understanding voice health. Through direct interviews and surveys with teachers and other professionals, we gain insight into their lived vocal experiences, workplace challenges, and individual adaptation strategies. This person-centered approach ensures that assessment and intervention priorities align with the realities of occupational voice demands and individual needs.

This approach offers particular promise for high-risk occupational groups including teachers, call center workers, and professional voice users, where vocal demands often challenge sustainable workload thresholds. Implementation requires enhanced measurement protocols that distinguish biological from artifactual variability, repeated functional assessments to establish individual baselines, and tracking of individual adaptation patterns in response to vocal loading.

KEYWORDS: precision medicine; vocal variability; individualized assessment

PRESENTER BIO: Mark Berardi is a Research Scientist whose academic journey began in acoustical physics with applications in music, speech, and architectural acoustics. He earned his PhD from Michigan State University developing models for aging voice, vocal effort, and vocal fatigue. He completed postdoctoral fellowship at University Hospital Bonn, investigating neurobiological factors in vocal production.

ACKNOWLEDGMENTS: This research was supported by the NIDCD of the NIH under Grant No. R01DC012315. The content is solely the responsibility of the authors and does not necessarily represent the official views of the National Institutes of Health.

DISCLOSURES: None of the authors have conflicts of interest to disclose.

REFERENCES

- Berardi, M. L., Whitling, S., & Hunter, E. J. (2025). Voice fatigue subtyping through individual modeling of vocal demand responses. *Scientific Reports*, 15(1), 1–11. <https://doi.org/10.1038/s41598-025-10565-2>
- Hunter, E. J., & Berardi, M. L. (In Press). Embracing variability: Toward proactive and precision-based voice science. *Logopedics Phoniatrics Vocology*. <https://doi.org/10.1080/14015439.2025.2562929>
- Hunter, E. J., Cantor-Cutiva, L. C., & Walden, P. R. (In Press). From Screening to Precision: Searching for Voice Disorder-Specific Acoustic and Auditory-Perceptual Metrics. *Journal of Voice*. <https://doi.org/10.1016/j.jvoice.2025.09.007>