

1. This document is created for accessibility and is intended to be a supplemental resource. If you would like to print a hard copy of this document, please follow the general instructions below to print multiple slides on a single page or in black and white.
2. This handout is for reference only. Non-essential images have been removed for your convenience. Any links included in the handout are current at the time of the live webinar but are subject to change and may not be current at a later date.
3. Copyright: Images used in this course are used in compliance with copyright laws and where required, permission has been secured to use the images in this course. All use of these images outside of this course may be in violation of copyright laws and is strictly prohibited.
4. Social Workers: For additional information regarding standards and indicators for cultural competence, please review the NASW resource: Standards and Indicators for **Cultural Competence in Social Work Practice**
5. Need Help? Select the “Help” option in the member dashboard to access FAQs or contact us.

How to print Handouts

On a Mac

- Open PDF in Preview
- Click File
- Click Print
- Click dropdown menu on the right “preview”
- Click layout
- Choose # of pages per sheet from dropdown menu
- Checkmark Black & White if wanted.

How to print Handouts

On a PC

- Open PDF
- Click Print
- Choose # of pages per sheet from dropdown menu
- Choose Black and White from “Color” dropdown

No part of the materials available through the continued.com site may be copied, photocopied, reproduced, translated or reduced to any electronic medium or machine-readable form, in whole or in part, without prior written consent of continued.com, LLC. Any other reproduction in any form without such written permission is prohibited. All materials contained on this site are protected by United States copyright law and may not be reproduced, distributed, transmitted, displayed, published or broadcast without the prior written permission of continued.com, LLC. Users must not access or use for any commercial purposes any part of the site or any services or materials available through the site.

Maturation of Speech Recognition for Noisy Environments, in partnership with American Auditory Society

Lori Leibold, PhD

Senior Director of Hearing Research
Boys Town National Research Hospital



This course may not be copied, reproduced, published, displayed, broadcast, reduced to any electronic medium or machine-readable form, or otherwise used or distributed in any form without the prior written consent of continued.com LLC. © 2025 Continued®

Disclosures

- **Presenter Disclosure:** Financial disclosures: Dr. Leibold has grant support provided by the National Institutes of Health (NIH) and is an employee of Boys Town National Research Hospital. Non-financial disclosures: Dr. Leibold is an ad hoc reviewer for the NIH and the Hearing Health Foundation; and a member of the American Auditory Society, the Acoustical Society of America, and the Association for Research in Otolaryngology.
- **Content Disclosure:** This learning event does not focus exclusively on any specific product or service.
- **Sponsor Disclosure:** This course is presented in partnership with American Auditory Society.

Limitations/Risks

- The applicability and accuracy of the course content may vary by the clinical, geographical, and cultural context; you should evaluate the course accordingly.
- As healthcare is rapidly evolving, new findings may change the validity of the information presented in this course. Providers should always consult the latest research and guidelines from professional associations.
- The interventions discussed in this course may be limited in generalizability, requiring practitioners to consider specific individual and population factors.
- The course may not cover all possible risk factors, diagnostic methods, or treatment options, and complex cases may require additional considerations.
- Healthcare providers should apply cultural competence and sensitivity to ensure effective and respectful care based on their patients' diverse needs.
- Providers should always evaluate the limitations of diagnostic and screening tools, considering their potential for false results as well as any ethical implications.
- It is the responsibility of course participants to critically evaluate the evidence from multiple sources to guide professional practice.

Learning Outcomes

After this course, participants will be able to:

- Describe how speech recognition in the presence of competing background sounds develops during infancy and childhood.
- Discuss how factors such as age, language skills, and hearing loss may impact the ability to hear speech in noise.
- Propose potential clinical tools for accurately assessing speech recognition in steady noise and in competing speech.

Collaborators and Support



Emily Buss



Lauren Calandruccio



Nicole Corbin



Heather Porter



Ryan McCreery



Christina Dubas



THE UNIVERSITY
of NORTH CAROLINA
at CHAPEL HILL

How does the ability to recognize speech in the presence of competing sounds develop?



How do we balance scientific rigor with real-world validity?



Presentation Outline

Conceptual
Framework &
Background

Key Scientific
Findings

Clinical
Implications &
Application

Unresolved
Questions

Summary &
Future Directions

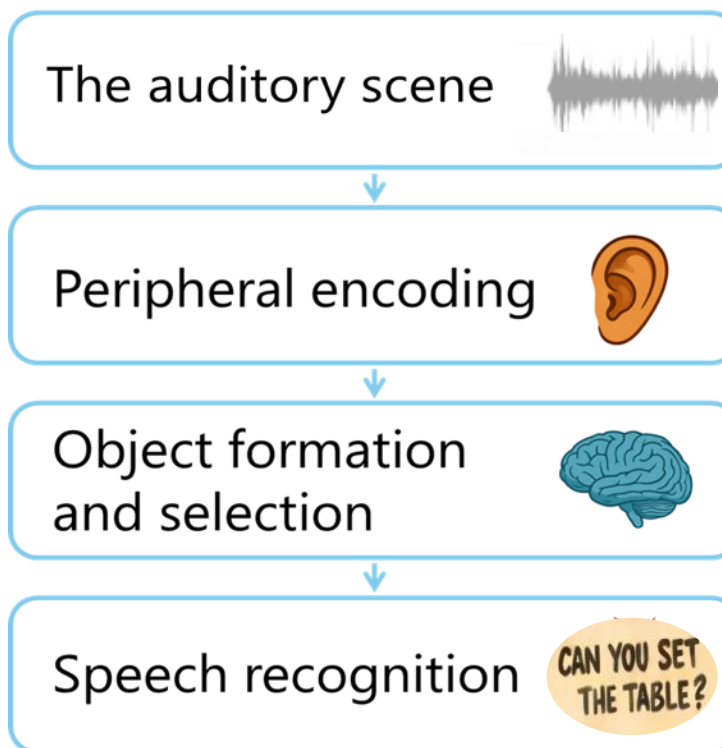


What are we trying to understand?



Stages of Auditory Processing

Immaturity within any stage of processing can influence how well children perceive speech in their everyday lives.



The auditory scene is the combination of all sounds that reach a listener.



Masker



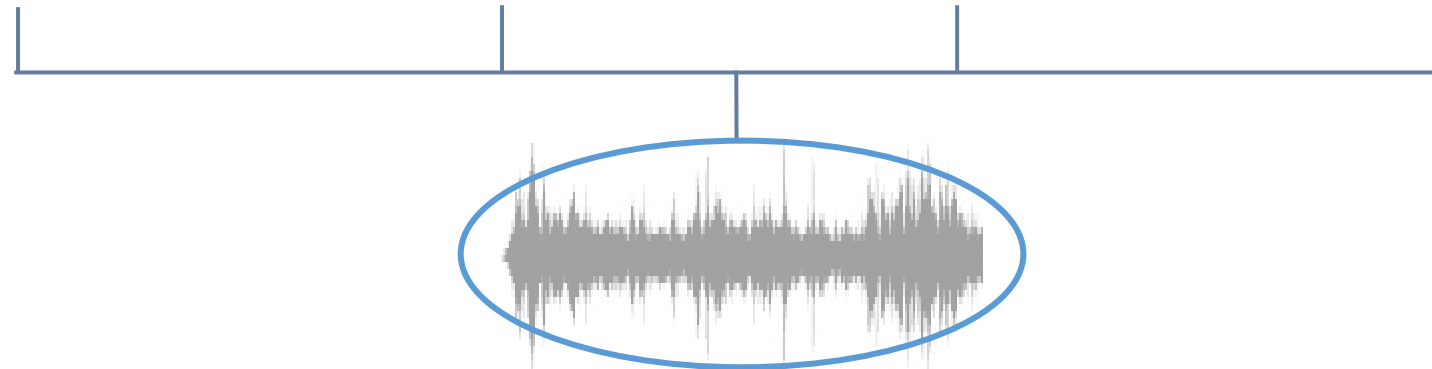
Target



Masker

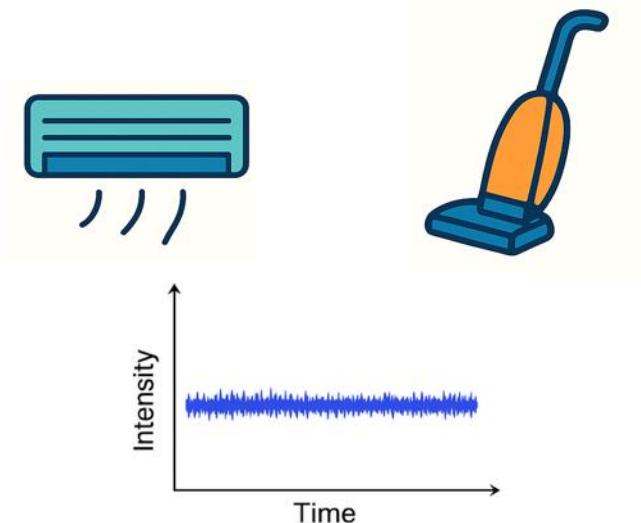


Masker

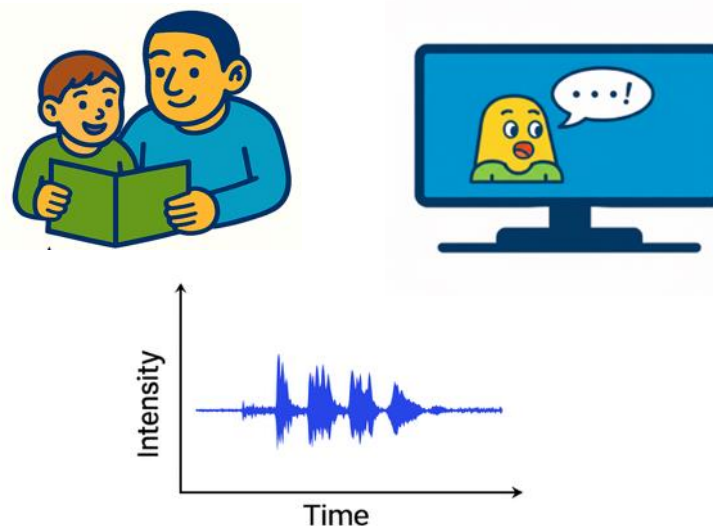


Children's auditory scenes are composed of multiple sources of competing sounds.

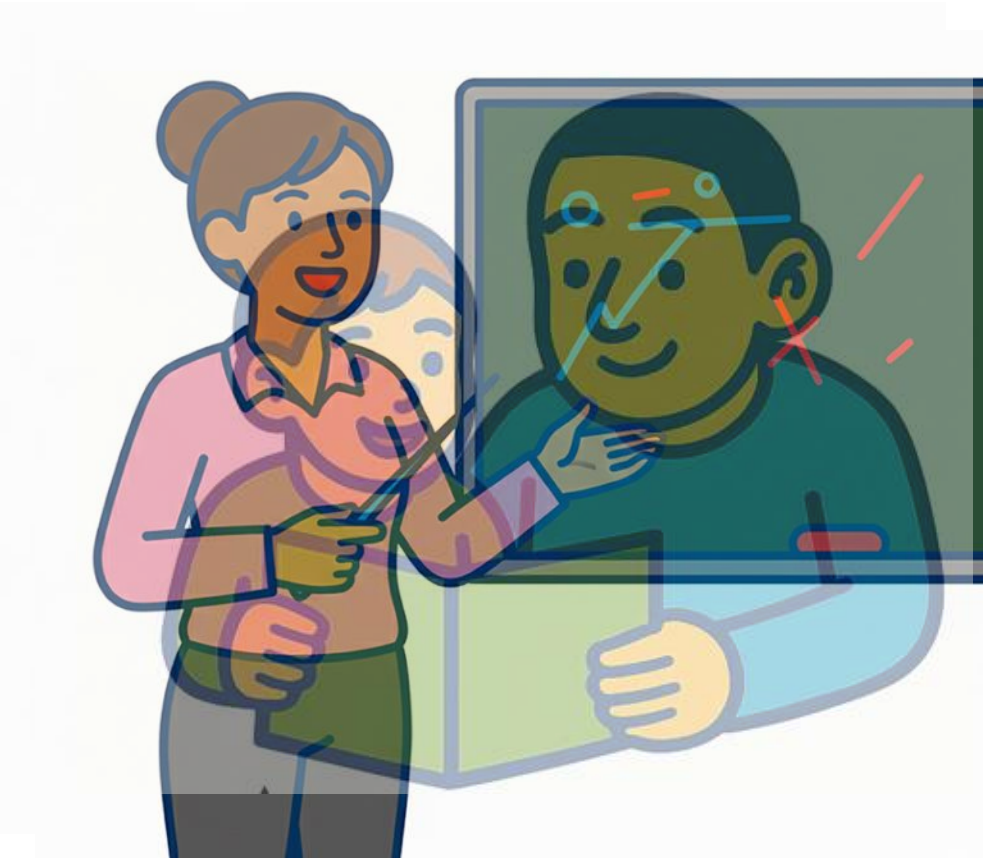
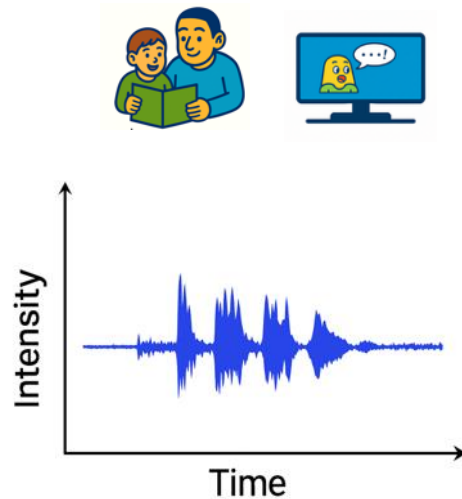
Some sources produce relatively steady sound.



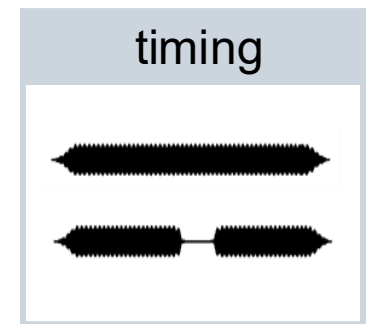
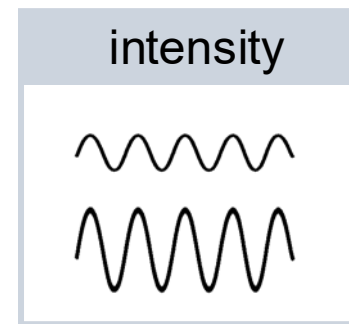
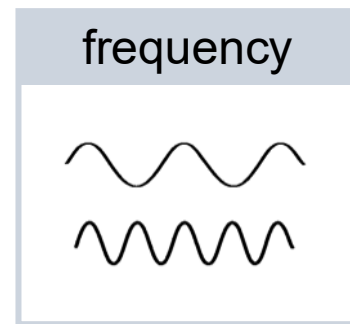
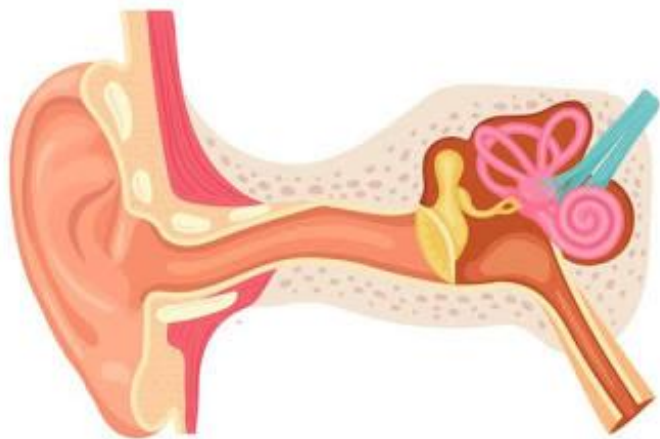
Other sources produce more dynamic sounds.



Different processes are responsible for masking with noise and competing speech.



The spectral, intensity, and temporal characteristics of the sound mixture must be adequately represented by the peripheral auditory system.



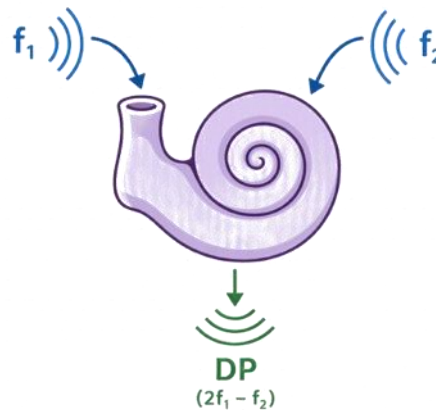
Peripheral encoding reaches adult-like resolution by 6 months of age.

Anatomical studies



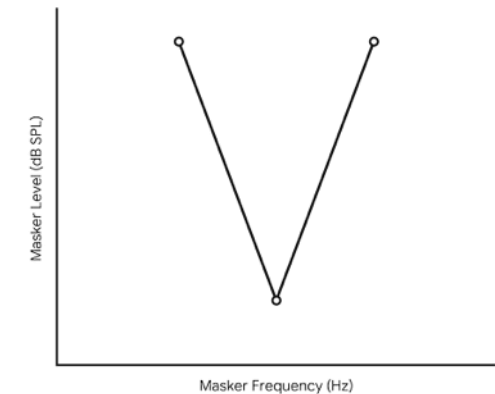
Rask-Andersen et al., 2012

DPOAE tuning curves



Abdala & Keefe, 2011

Psychophysical tuning curves



Spetner & Olsho, 1990

#1: Children are worse at recognizing masked speech than adults, particularly when the masker is speech.



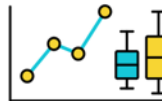
Children and adults with normal hearing

- 5-16 years; n=56
- 18-44 years; n=16

Adaptive, open-set procedure

Target speech: monosyllabic words

Masker conditions: (1) speech-shaped noise
(2) two streams of competing speech

Results 

SRTs for 5- to 7-year-olds were **~4 dB** worse than 13- to 16-year-olds and adults in noise.

This age effect was **~7 dB** in competing speech.

#2: Speech-in-speech recognition is one of the last auditory skills to develop.



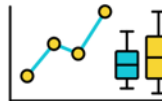
Children and adults with normal hearing

- 5-16 years; n=33
- 18-35 years; n=10

Adaptive, closed-set, 4AFC procedure

Target speech: disyllabic words

Masker: one stream of competing speech

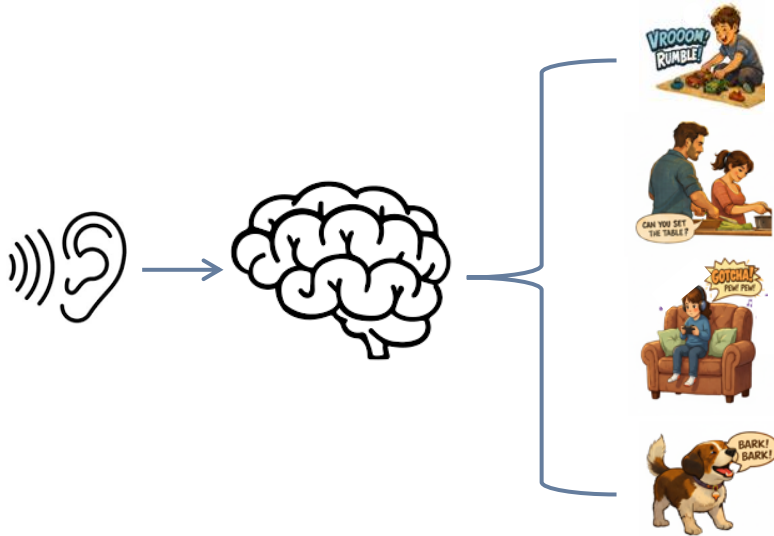
Results 

SRTs improved by more than **16 dB** during childhood.

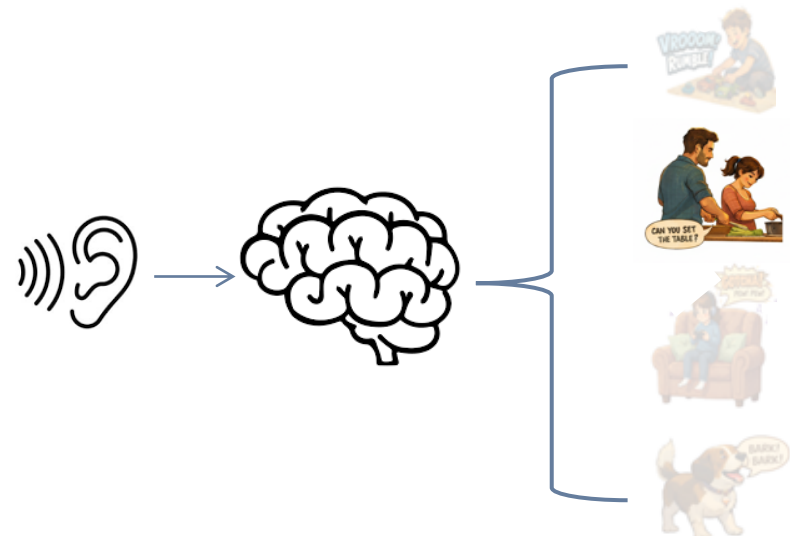
Mature performance was not observed until adolescence (**~13 years of age**).

#3: Speech-in-speech recognition is influenced by central auditory processing.

Grouping sounds into separate auditory objects



Selectively attending to target speech



School-age children, like adults, benefit from a mismatch in target/masker sex.



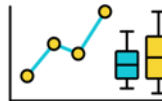
Children and adults with normal hearing

- 5-10 years; n=10
- 19-28 years; n=8

Adaptive, closed-set, 4AFC procedure

Target speech: disyllabic words produced by a male (matched) or female (mismatched)

Masker: two streams of speech, each produced by a male

Results 

Children performed more poorly than adults overall.

Children and adults benefitted by a similar amount when target and masker speech were mismatched in sex relative to when they were the same sex.

Infants do not appear to benefit from a mismatch in target/masker sex.



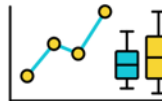
Infants and adults with normal hearing

- 7-13 months; n=18
- 18-26 years; n=18

Adaptive, observer-based procedure

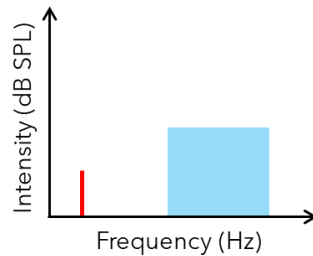
Target speech: a disyllabic word produced by a male (matched) or female (mismatched)

Masker: two-male-talker speech

Results 

Using the same stimuli, infants performed similarly when the target and masker speech were produced by talkers mismatched in sex and when talkers were the same sex.

Infants and young children listen less selectively than adults.



Children and adults with normal hearing

- 4-9 years; n=23
- 19-33 years; n=8

Adaptive, 2IFC detection task

Signal: 1 kHz tone

Conditions: quiet, remote noise (4-10 kHz)

Amount of Masking = $SRT_{noise} - SRT_{quiet}$



4- to 6-year-olds have elevated detection thresholds for a 1000 Hz tone when it is presented with a remote-frequency noise than when it is presented in quiet.

7- to 9-year-olds and adults show no masking with the remote-frequency noise.

#4: Immature language skills impact children's speech-in-speech recognition.



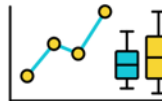
Preschoolers with normal hearing

- 3-4 years; n=23

Adaptive, closed-set, 4AFC procedure with
'active' familiarization and training

Target speech: disyllabic words

Masker: two streams of competing speech

Results 

Preschoolers with larger vocabularies had lower
(better) SRTs for speech-in-speech recognition
than preschoolers with smaller vocabularies.



Children with hearing loss struggle more than peers with normal hearing in noisy environments.

Peripheral encoding



Hearing loss impacts the fidelity with which sounds are represented by the peripheral auditory system.

Object formation
and selection



Children with hearing loss may have reduced, degraded, and/or variable experience with sound.

Hearing loss exacerbates children's speech-in-noise difficulties.



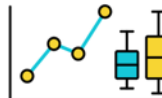
Children with moderate to severe hearing loss tested wearing amplification and children with normal hearing

- Hearing loss: 9-17 years; n=17
- Normal hearing: 9-11 years; n=10

Adaptive, closed-set, 4AFC procedure

Target speech: disyllabic words

Masker: speech-shaped noise

Results 

On average, children with hearing loss required an SNR that was **3.5 dB** higher than peers with normal hearing to recognize speech in noise.

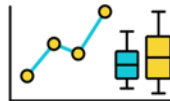
SRTs measured in the noise masker were not associated with parent reports of their child's listening difficulties.

Children with hearing loss are at an even greater disadvantage when multiple people are talking.



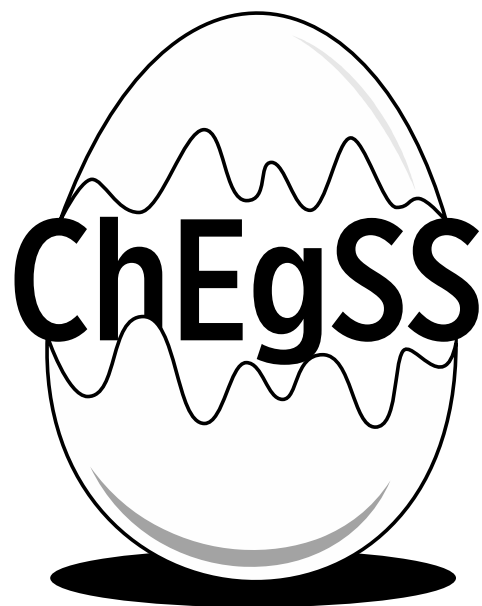
Target speech: disyllabic words

Masker: two streams of competing speech

Results 

The disadvantage for children with hearing loss in the speech masker was **8.1 dB**.

Speech-in-speech recognition performance was associated with parent reports of their child's listening difficulties.



Children's English and Spanish Speech Recognition Test



Carhart



Hall



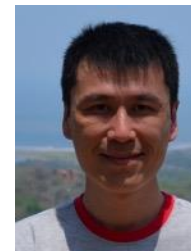
Elliot



Gravel



Stelmachowicz



Shi

What is the ChEgSS Test?



A computer-based tool for assessing closed-set word recognition in English and in Spanish, in quiet or with a masker that is either speech-shaped noise or competing speech.



The tester does not need to speak the test language or score productions in that language.



ChEgSS uses an adaptive, forced-choice procedure with a picture-pointing response.

Targets: 30 disyllabic words in English and Spanish

- Spoken by the same female talker
- Appropriate vocabulary for a 5-year-old
- Easily illustrated

Masker Conditions: (1) speech-shaped noise (2) two streams of English speech (3) two streams of Spanish speech

feather/pluma



woman/mujer



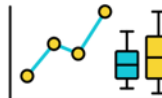
zebra/cebra



Extensive normative data were collected on 4- to 18-year-old children.



Participants: 83 Spanish-English bilingual and 89 English monolingual children with normal hearing
Bilingual children completed testing in English and in Spanish, in the noise and the speech masker
Monolingual children completed testing twice in English

Results 

Performance was similar in both languages for bilinguals with adequate proficiency.

Prediction intervals have been constructed, for comparison with clinical data.

Are we capturing the challenges children face in their everyday lives?



Some of our tasks are likely too difficult (relative to the real world).



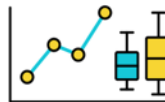
Children and adults with normal hearing

- 5-14 years; n=33
- 18-39 years; n=16

Adaptive, open-set sentence recognition

Target speech: BKB sentences produced using a clear or a conversational style

Masker: two streams of speech

Results 

Young children did not show a clear speech benefit.

The answer is different when the task is easier (providing an additional binaural cue)



Children and adults with normal hearing

- 4-7 years; n=20
- 19-39 years; n=10

Adaptive, open-set sentence recognition

Target speech: BKB sentences produced using a clear or a conversational style

Masker: two streams of speech

Now with a binaural cue! $M_0 T_\pi$

Results 

Young children and adults showed a similar clear speech benefit.

Generalization is limited by sampling bias.



Considering the child population of the US:

- 16% live in poverty
- 21% speak a language other than English
- 15% receive special education services
- 20% live in rural communities

Speech-in-speech recognition is challenging for individuals with Down syndrome.



The Lifespan Study of Hearing and Balance in Individuals with Down Syndrome

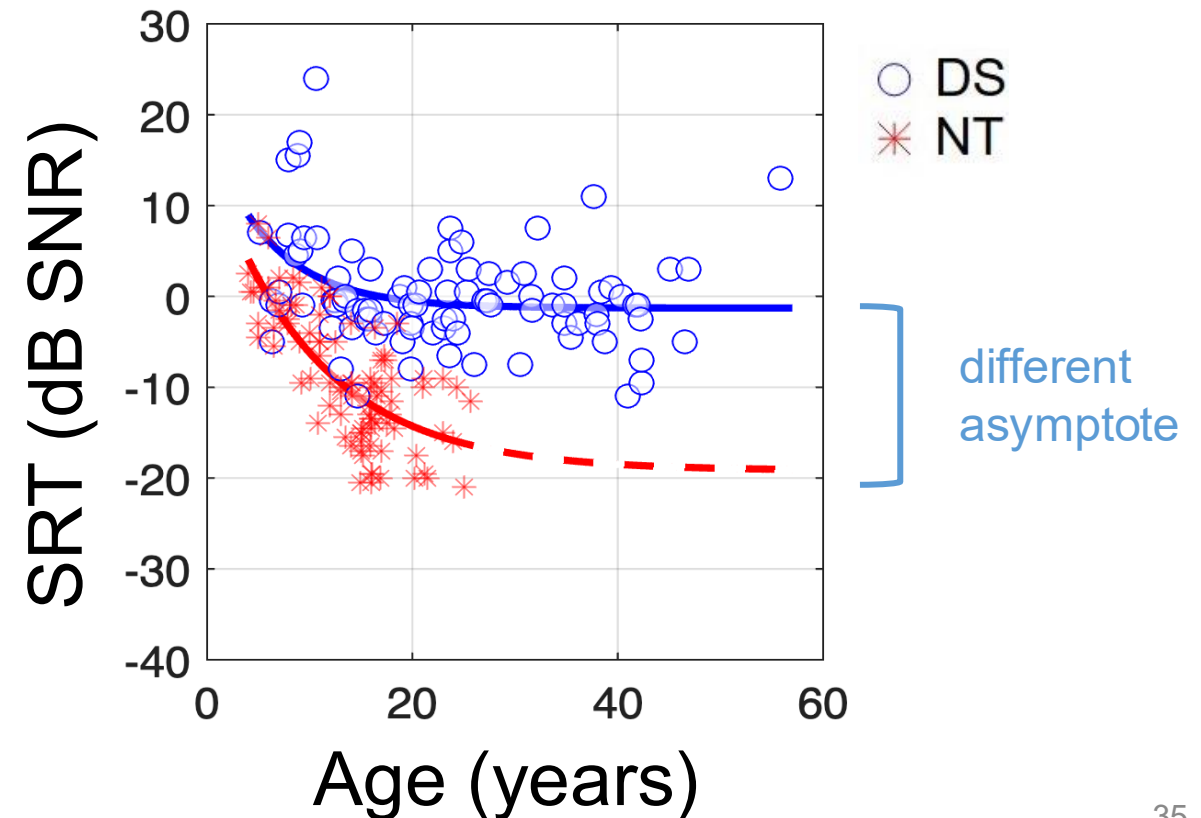
Individuals with Down syndrome and individuals who are neurotypical

- DS: 5-55 years; n=92
- NT: 4-25 years; n=90

Adaptive, closed-set, 4AFC procedure

Targets: disyllabic words

Masker: two streams of competing speech





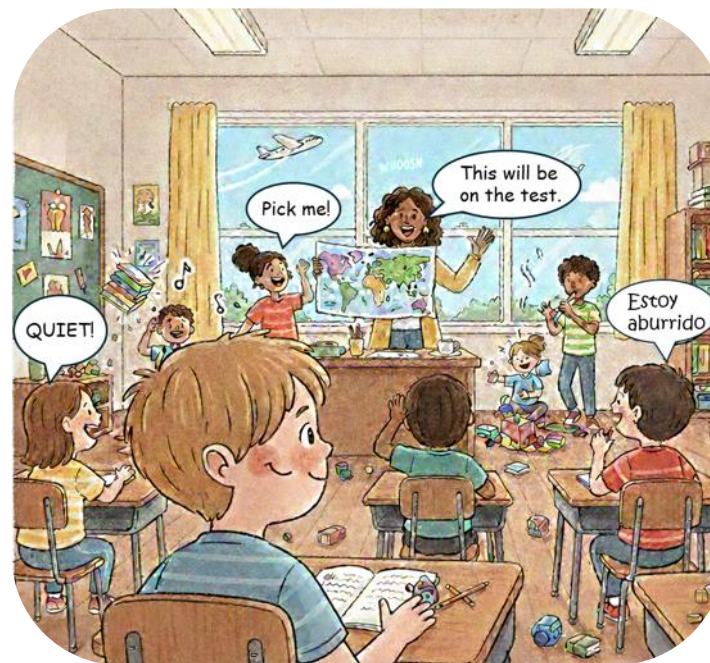
What a child hears in a noisy environment is not what an adult hears.



The ability to recognize speech when multiple people are talking at the same time requires years of experience and maturation of central auditory processing.

Consistent access to high-quality auditory input provides the foundation for language development, social functioning, and school success.

Not all children are afforded this access.



Our future directions include characterizing masked speech recognition for...

Speech produced by a wider range of talkers and styles



A community-based sample of children





Summary & Future Directions



Thank you!

Lori.Leibold@boystown.org

Congratulations!

- You've completed the course and can move on to the exam!
- From your account:
 - Go to pending Courses
 - Choose take exam