

Nomeação após instrução por Múltiplos Exemplares em crianças com deficiência auditiva

Naming following Multiple Exemplar Instruction in children with hearing loss

 ANA CLAUDIA M. ALMEIDA-VERDU^{1, 4}

 ALINE DANIELA G. S. VIEIRA^{2, 3}

 FERNANDA DE OLIVEIRA ORTIGOZA¹

 FELIPE MONTEIRO-SILVA¹

 MARIA CLARA L. BRISENO¹

 SARAH A. LECHAGO³

¹SÃO PAULO STATE UNIVERSITY

²UNIVERSITY OF SÃO PAULO

³UNIVERSITY OF HUSTON-CLEAR LAKE (USA)

⁴NATIONAL INSTITUTE OF SCIENCE AND TECHNOLOGY ON BEHAVIOR

Abstract

Auditory Neuropathy Spectrum Disorder (ANSD) is a type of hearing impairment wherein people can detect auditory stimuli, but sounds are processed in a desynchronized manner. This affects sound discrimination, dictated word recognition, pictures and events naming, and, therefore, vocabulary acquisition and expansion. Cochlear Implants (CI) are used as an intervention approach that enables sound detection, but they require auditory skills learning. Multiple Exemplar Instruction (MEI) has been used to promote the integration between listening and speaker behaviors, as well as the emergence of naming in children with minimal verbal repertoires. Speaker unidirectional naming is a subtype in which listening responses to specific items are directly taught, and, following the instruction, the child can respond as a speaker to the same items, naming them, without direct training. This study evaluated MEI on the integration between listening and speaker behaviors, as well as on the emergence of speaker unidirectional naming in two boys, aged 6 and 10, diagnosed with ANSD and users of CIs. Selection-based listening behavior, echoic responses, and tacting pictures were included in MEI and assessed pre- and post-intervention, with three sets of stimuli. The third set of stimuli evaluated the emergence of a new repertoire after exposure to MEI and the results were positive in unidirectional speaker naming for both participants. The results are discussed in terms of the experimental control for future research and the potential clinical application of MEI in the auditory rehabilitation of children with ANSD and CIs, as well as its relevance to future research directions.

Palavras-chave: verbal behavior, hearing impairment, Multiple Exemplar instruction, naming.

Resumo

O Espectro da Neuropatia Auditiva (ENA) é um tipo de deficiência auditiva em que a pessoa detecta estímulos auditivos, mas os sons são processados de maneira dessincronizada. Isso afeta a discriminação sonora, o reconhecimento de palavras ditadas, a nomeação de figuras e eventos e, consequentemente, a aquisição e a ampliação do vocabulário. O implante coclear (IC) é utilizado como uma abordagem de intervenção que possibilita a detecção sonora, mas requer a aprendizagem de habilidades auditivas. A Instrução por Múltiplos Exemplares (MEI) tem sido utilizada para promover a integração entre comportamentos de ouvinte e de falante, bem como a emergência de nomeação em crianças com repertórios verbais mínimos. A nomeação unidirecional de falante é um subtipo em que respostas de ouvinte a itens específicos são ensinadas diretamente e, após a instrução, a criança pode responder como falante aos mesmos itens, nomeando-os, sem treino direto. Este estudo avaliou os efeitos da MEI sobre a integração entre comportamentos de ouvinte e de falante, bem como sobre a emergência de nomeação unidirecional de falante, em dois meninos, com 6 e 10 anos de idade, diagnosticados com ENA e usuários de IC. O comportamento de ouvinte baseado em seleção, as respostas ecoicas e os tatos de figuras foram incluídos na MEI e avaliados antes e depois da intervenção, com três conjuntos de estímulos. O terceiro conjunto de estímulos avaliou a emergência de um novo repertório após a exposição à MEI, e os resultados foram positivos para a nomeação unidirecional de falante em ambos os participantes. Os resultados são discutidos em termos do controle experimental para pesquisas futuras e da aplicação clínica potencial da MEI na reabilitação auditiva de crianças com ENA e IC, bem como de sua relevância para direções futuras de pesquisa.

Keywords: comportamento verbal, deficiência auditiva, Instrução por Múltiplos Exemplares, nomeação.

NOTES. WE THANK BIANCA AUGUSTO AND TAFNES IKEGAMI PEREIRA, UNDERGRADUATE PSYCHOLOGY STUDENTS, FOR DATA COLLECTION UNDER THE SUPERVISION OF THE FIRST AUTHOR. THIS WORK WAS CONDUCTED AS PART OF THE SCIENTIFIC PROGRAM OF THE NATIONAL INSTITUTE OF SCIENCE AND TECHNOLOGY ON BEHAVIOR, COGNITION AND TEACHING (INCT-ECCE), FUNDED BY SAO PAULO RESEARCH FOUNDATION (FAPESP; GRANT # 2014/50909-8), THE NATIONAL COUNCIL FOR SCIENTIFIC AND TECHNOLOGICAL DEVELOPMENT (CNPQ; GRANT # 465686/2014-1), AND CAPES (GRANT # 88887136407/2017-00). ALINE DANIELA G. S. VIEIRA WAS SUPPORTED BY A CAPES MASTER'S SCHOLARSHIP DURING THE WRITING OF THE MANUSCRIPT (GRANT #88887.670929/2022-00) AND BY CAPES (GRANT #88887.929451/2023-00) AND FAPESP (GRANT #2024/10070-0) DOCTORAL SCHOLARSHIPS DURING THE SUBMISSION AND REVISION PROCESS. ANA CLAUDIA MOREIRA ALMEIDA-VERDU HAS RESEARCH PRODUCTIVITY GRANTS FROM CNPQ (GRANTS # 305522/2021-3). WE THANK THE CHILDREN AND CAREGIVERS FOR THEIR PARTICIPATION

✉ ana.verdu@unesp.br

DOI: [HTTP://DX.DOI.ORG/10.18542/REBAC.V22I1.20848](http://dx.doi.org/10.18542/REBAC.V22I1.20848)

Auditory Neuropathy Spectrum Disorder (ANSD) is a type of hearing loss which affects the cochlear nerve and leads to a desynchrony in sound stimuli conduction. It can be either a unilateral or bilateral auditory loss, which is not necessarily symmetrical and ranges from mild to profound. In individuals with ANSD, the variability of hearing thresholds and impairment in sound processing might lead to difficulties in recognition and speech production. This variability in auditory-verbal performance can be challenging for verbal vocal communication learning. (Costa et al., 2012; Norrix & Velenosky, 2014).

The cochlear implant (CI) is an electronic device that might be a viable alternative, due to its benefits in auditory recognition learning and speech acquisition in people who have severe hearing loss. The CI directly stimulates the cochlear nerve, partially replacing the sensorial cells functions and, in cases of ANSD, it might benefit neural synchrony (Carvalho et al., 2011; Fernandes et al., 2015). Although the CI reestablishes ambient sounds detection, establishing other auditory (i.e., discrimination, recognition, and comprehension) and verbal (i.e., speech production with accuracy) skills requires learning (Almeida Verdu et al., 2014; Fernandes et al., 2015).

Most children who do not have hearing loss learn new vocal verbal language and communication skills with relative ease, by observing others in the natural environment. For example, a child may observe an individual point to a remote control and ask another individual, "Can you please hand me the remote control?", which is followed by that second individual retrieving the device and handing it to the requestor. Echoic responding is implicated in this scenario wherein the individual engaged in listener behavior by looking at the remote and echoing the original requestor when they ask for the remote. In a new circumstance, without direct training, most children would be able to respond as both speaker and listener by pointing to the remote control if someone asked for it or tact it if someone were to ask, "What is it?". This scenario, demonstrated by several studies, involves an integration between speaker and listening behaviors (Greer & Speckman, 2009; Santos & Souza, 2020) if the learners consistently engage in the echoic while looking at the item and listening to the original requestor mand for it. This integration depends on the establishment of two other behaviors, listening and tacting, and the repeated and integrated occurrence of the echoic, tact, and listener repertoires (Horne & Lowe, 1996; Kruger et al., 2024; Santos & Souza, 2020).

This study targets two types of speaker behavior, objects labeling (i.e., tact), repeating words (i.e., echoic) and listener behavior. Tact is a verbal operant described by the emission of responses evoked by non-verbal stimuli or their properties, whether they are public (i.e., objects, events, pictures) or private non-verbal stimuli (i.e., an organism's internal states and feelings). This response is maintained by social reinforcement because it enables listeners to expand their contact with the environment. It is in interaction with another person that conditions are created for people to observe events and respond vocally by naming them (Skinner, 1957/1978). Tacting may have many extensions beyond its basic property (i.e., metaphors and abstractions), with its boundaries defined by the verbal community.

The echoic is presented as early as the initial stage of typical development, when babbling becomes differentiated and under the influence of the interlocutor's speech. It must have point-to-point correspondence between the auditory discriminative stimulus (speech) and the echoic response (repetition) and must be issued right after the model is presented, and followed by social reinforcement (Skinner, 1957/1978).

The echoic and the tact may require the same vocal topography, as saying "cake" after an adult says "cake" (echoic) and saying "cake" when seeing a cake (tact), but they are functionally independent, and the acquisition of one does not guarantee acquisition of the other. Listener behavior is described as a physical non-verbal response under the influence of a verbal stimulus. In the example with the cake, if an adult says, "who wants cake?" a child playing in

another room might turn their head toward the sound source. Listener (seeing or pointing to what others say) and speaker behaviors (naming objects and repeating what others say) are also functionally independent (Guess, 1969; Hussein et al., 2018) suggesting that the acquisition of one does not guarantee the acquisition of the other in the absence of direct training.

The integration between listener and speaker behaviors might not be learned in naturalistic conditions by many children with a minimal verbal repertoire, and children with ANSD are at particular risk of not acquiring it. Therefore, procedures and programs that can establish and enhance such skills are implicated for those individuals with language deficits (Greer & Ross, 2008), including children with ANSD with CI who, following implementation of the device, are learning to discriminate auditory stimuli, relate them to everyday events, and produce corresponding vocalizations (Almeida-Verdu et al., 2014). In light of this, developing and evaluating instructional procedures that promote the establishment and expansion of such skills is warranted.

Research on the integration between hearing and speaking has both scientific and social relevance to individuals who have a history of language delay or hearing loss. Furthermore, children with hearing loss may exhibit variability in vocal responses, lacking point-to-point correspondence, which reduces speech intelligibility (Yoder et al., 2005). This may consequently restrict opportunities for social interaction, incidental learning, and contact with reinforcement contingencies mediated by the verbal community, resulting in significant impairments. Furthermore, imprecision and variation in speech accuracy led to impairments in social interactions, such as poor intelligibility or lack of comprehension of speech by the listener, thus compromising access to reinforcer mediated by others.

Recent research has evaluated procedures to produce functional interdependence or integration between listening and speaker behaviors (Almeida Verdu et al., 2014; Pereira et al., 2016; Rique et al., 2017) and speech accuracy in children with auditory disabilities and CI. One such strategy is Multiple Exemplar Instruction (MEI), which consists of teaching a subset of responses under a variety of different speaker and listener contexts (e.g., pointing-to-picture; labelling; requesting; repeating). After continuing repetition of these tasks' sets, with different response topographies (pointing, vocalizing) and different stimuli control (dictated words, pictures), there might be an emergence of speaker responses following the listening training (and vice versa) or following incidental experiences. For example, after learning how to listen to a dictated word and point to the corresponding picture, a child might be able to label such a picture without direct teaching. This type of naming is designated as unidirectional speaker naming, described by emergent speaker behavior resulting from direct training of listener behavior; that is, the listener behavior is trained, and the speaker behavior is tested. Unidirectional speakers differ from bidirectional naming in which both listener and speaker behavior emerge, without a history of direct training. (Hawkins et al., 2009; 2018).

In a systematic review, 18 studies that evaluated the emergence of unidirectional naming in speaker responses were found (Santos & Souza, 2020). The studies included typically developing participants (i.e., Bandini et al., 2012; Gilic & Greer, 2011; Miguel et al., 2008); children and adolescents with Autism Spectrum Disorder (ASD) diagnostic (i.e., Hawkins et al., 2009; Olaff et al., 2017; Santos & Souza, 2016); both, typically developing children, and children with ASD (i.e., Longano & Greer, 2015); children with developmental delays (i.e., Greer et al., 2005); and children and adolescents with auditory disabilities and CI (i.e., Pereira et al., 2016; Pereira et al., 2018). Some authors point to MEI as a promising procedure, whose clinical potential must be further explored to increase the efficiency of language programs for children who have been diagnosed with developmental disabilities (Grow & Kodak, 2010). To verify whether MEI, as an independent variable, produces any effects on the emergence of speaker unidirectional naming, as a dependent variable, this condition (i.e., listening behavior and speaker test) must be assessed before and after the exposure to MEI (Santos & Souza, 2020). The MEI procedure has been evaluated with children with cochlear implants and ANSD (Merlin et al., 2019; Rique et al. 2017). Before the teaching, listening, and speaker behaviors presented very disparate percentages of correct responses; after the teaching, the correct response percentages tend to overlap. However, the necessary tests for evaluating speaker unidirectional naming (Hawkins et al., 2018) were conducted by Mascotti et al. (2024). In their study, three children with ASD and only one child with ANSD and cochlear implant were exposed to MEI, and tact emergence was demonstrated following the listening training. For these studies, MEI consisted of selection-based listening behavior, as well as echoic and tact behaviors. The same vocal behavior that is emitted in the echoic task is emitted in the picture naming task. And the same stimuli are involved in both echoic and selection-based listening behavior.

Although previous studies have demonstrated the effectiveness of MEI on the integration of speaker and listener repertoires (Mascotti et al., 2024; Merlin et al., 2019; Rique et al., 2017), there are gaps in the literature. For instance, Pereira et al. (2018) and Mascotti et al. (2024) included a third set of stimuli, totally new, which trained listening behavior and tested speaking behavior following MEI experience. This stage was critical to demonstrate the role of MEI

in the integration between listening and speaker behavior. In other previous research on MEI with children with hearing disabilities, this aspect had not been measured yet (Merlin, et al., 2019; Pereira et al., 2016; Rique et al., 2017). In Mascotti et al., only one participant had ANSD, and MEI was composed of echoic, tact, and listening behaviors. As for Pereira et al., the type of hearing disability was not specified, and MEI did not include echoic behavior. Although these studies indicate the potential of MEI to promote integration between listener and speaker repertoires, evidence on this effect specifically in children with ANSD is still limited, especially regarding the adequate assessment of the MEI's inducing role on the speaker's unidirectional naming. Adequate assessment consists of exposure to listener training and speaker testing with a third set of stimuli that did not participate in the demonstration of functional independence of hearing and speaking (Set 1) and that did not participate in MEI (Set 2).

The purpose of the present study was to systematically replicate the results of Mascotti et al. (2024), by evaluating the effects of an MEI arrangement composed of listener, echoic, and tact repertoires on the integration between listener and speaker behaviors, as well as on the emergence of unidirectional speaker naming in children with ANSD using cochlear implants, in an applied context.

Method

Participants

The participants were two boys, John and Harry (pseudonyms), who were receiving treatment at a university psychology clinic. The demographic and diagnostic data were collected through an interview. John was six years old and had ANSD diagnosis. He had been using bilateral cochlear implants for five years and five months. Harry was 10 years old. In addition to an ANSD diagnosis, Harry had difficulties with orofacial motricity skills, that is, difficulties involving the coordination and movements of the mouth, face, and neck, which are relevant to speech production. He had been using a unilateral cochlear implant in his right ear for five years and six months. In his other ear, he had a hearing aid device. The data were collected through caregiver interviews, without access to speech-language pathology records. So, the correct age at diagnosis and the use time of CI wasn't available. The responses for listening and speaker behaviors were pre-tested. The highest results, for both, were in echoic behavior, listening behavior presented variability, and tact behavior had the lowest results. These results are described in greater detail in the "Results" section.

Settings and Variables

For Harry, sessions were conducted in a school clinic of a psychology faculty in the countryside of the state of São Paulo, and John's session were conducted at home. The activities were carried out in a well-lit room; the environment was airy, and there was no co-occurring noise.

The primary dependent variable was the speaker's unidirectional naming which consisted of increasing the percentage of correct tact, after listener training with the corresponding stimuli. This means the participant was trained to engage in a listener response (point to a stimulus) and the researchers probed for the emergence of tact for the item depicted in the picture (speaker behavior), as well as for the echoic response (Hawkins et al., 2018), because it is a pre-requisite; both behaviors were monitored in the teaching and testing conditions. The analysis regarding tact and echoic responses vocalizations considered the participant's speech precision when compared with a typical listener. Echoic and tact responses were transcribed, and the point-to-point correspondence percentage was calculated, as detailed in the results analysis procedure section.

The primary independent variable for the MEI procedure involved selecting a subset of stimuli and teaching it under various contexts. Each stimulus in the set was rotated in three contexts: selection-based listening ("Point to"), tact ("What is this?"), and echoic ("Repeat") tasks.

Materials and stimuli selection

The stimuli used in this study were the same ones presented in the LAPIDAR (Mascotti & Almeida Verdu, 2020) program, structured according to the MEI teaching proposal, that is, the pictures were presented using PowerPoint® software, with a laptop; the dictated words used in listening and selection tasks were delivered by the researcher, following a trial sequence script. Files containing the stimuli and tasks in the PowerPoint® software can be accessed on the website <https://sgcd.fc.unesp.br/#!/lads/producao/produto-tecnico-e-tecnologico-programas/lapidar/>.

The illustrations represent easy-to-recognize dictated words, composed of two syllables and four phonemes each. There were multiple differences among the words (Guerra & Almeida-Verdu, 2020). After an initial assessment,

described in greater detail below, three sets composed of three stimuli (a total of nine stimuli) were selected for each participant. Figure 1 displays the three sets adopted for each participant.

Figure 1
Sets of Stimuli (Pictures and Words Dictated, and Their Corresponding Version in English) per Participant

	SET 1	SET 2	SET 3		SET 1	SET 2	SET 3
John	"Pata" (Paw) 	"Mapa" (Map) 	"Pena" (Feather) 	Harry	✓ "Taco" (Bat) 	"Tela" (Canvas) 	"Tico" (Tico) 
	"Taco" (Bat) 	"Pote" (Bowl) 	"Nabo" (Turnip) 		"Dora" (Dora) 	"Davi" (Davi) 	"Dama" (Checkers) 
	"Neve" (Snow) 	"Dora" (Dora) 	"Dama" (Checkers) 		"Pera" (Pear) 	"Pote" (Bowl) 	"Pata" (Paw) 

Experimental Design

The study used a single-subject experimental design with repeated probes across three sets of stimuli for each participant. Set 1 was selected to demonstrate the functional independence between listening training and tact emergence. After the MEI, the listening training was repeated with Set 1, followed by probes to see if tact emerged with the same set of stimuli. Finally, there was a new listening training, and a tact test was conducted with Set 3 as a control condition. Multiple tests were conducted throughout all phases of the procedure to monitor the conditions under which the tact would emerge and to control for threats to internal validity (Morin et al., 2024).

General Procedure

Three types of trials (three target responses) were adopted for each stimulus: listening behavior trials, picture tact trials, and word repetition trials. A computer screen was used to present the trials during each of the experimental conditions. Each trial included an instruction (S^D) which specified a response (pointing at or vocalizing), the participant's response, a programmed consequence for correct or incorrect responses, and a short interval between tasks. Listening training (pointing at the correct picture) started with three pictures displayed at the bottom of the screen, in central, left, and right positions. The experimenter issued the instruction, "Point to (name of the stimulus)". The correct response involved touching the picture that corresponded to the instruction; the participant had to respond within 5 s of the instruction. During training conditions, programmed consequences for correct or incorrect responses were provided. Correct responses were reinforced with social praise (e.g., "Very good!") and, when appropriate, with access to preferred items or activities. Incorrect responses were not followed by programmed consequences and resulted in the presentation of the next trial. Pointing to the picture that did not match the researcher's instruction, no responses with 5 s, or responses after a prompt were considered incorrect. The mastery criterion for the training phase was above 85% correct responses. There were no programmed consequences during pre- and post-test procedures.

For the echoic condition, a prompt for repetition trials was presented, such as a picture of a boy saying a word followed by the experimenter presenting the instruction (e.g., "Repeat [name of the stimulus]"). A correct response was defined as a response that repeated the word presented orally within 5 s of the instruction. During training conditions, there were programmed consequences for the correct or incorrect responses. Vocal responses were transcribed and analyzed by point-by-point correspondence with the target word; vocalizations with omissions, substitutions, or distortions were classified as incorrect or partially correct according to the phoneme score. The mastery criterion during the training phase was above 85% accuracy (point-to-point correspondence to the dictated word).

Tact probing involved presenting a picture in the middle of the screen; the experimenter asked, “What is that?” or “What’s the name of it?”. A correct response was defined as the participant saying the name that corresponded to the picture on the screen within 5s of the instruction. The vocal responses were transcribed and analyzed according to point-to-point correspondence with the vocalizations emitted from a listener; omissions, substitutions or distortions were classified as incorrect vocalizations. During training conditions, programmed consequences were provided for correct or incorrect responses, as well as echoic prompts for incorrect responses. The mastery criterion during the training phase was above 85% accuracy (point-to-point correspondence to the dictated word).

General Test of Stimuli

The general test of stimuli aimed to evaluate the general repertoire of the participants in three types of responses (i.e., listening, tact, and echoic). There were no programmed consequences for correct and incorrect responses. From these probes, the stimuli for which the child exhibited errors during the vocalization tasks were identified. These stimuli were organized into three sets, with three stimuli each, and were used in the subsequent phases of the procedure.

Pre-test. For the Pre-tests, three blocks, with nine trials each, were adopted to assess auditory recognition, word repetition, and picture naming. The three stimuli from each set (i.e., taco, pan, and turnip for Set 1) were tested three times in one block, with one trial for each of the target behaviors (i.e., “Point to taco,” “Repeat *taco*,” “What is this? (photo - taco)” and so on), for a total of 27 trials. There was no programmed consequence for correct or incorrect responses.

Listening training (pointing to) and speaker probes (echoic and tact) (Set 1). Stimuli from Set 1 were used for the listening training. The training started with the instruction “Point to (name)”. Pointing at the correct stimulus was followed by approval words, such as “Right!”. In case of incorrect responses, as a correction procedure, the instruction was repeated along with a gestural prompt, in which the experimenter pointed at the corresponding stimulus and then repeated the instruction. Each one of the blocks was structured with three stimuli from Set 1, in a total of nine trials, three for each stimulus. Totally correct responses or just one incorrect response in two consecutive trial blocks were considered as a learning criterion. Then, the participants were exposed to the next phase, the speaker tests. After the pointing response reached the mastery criterion for Set 1, both echoic and labeling responses, with Set 1, were probed. Teaching was terminated if the number of correct responses did not change for two consecutive sessions. Speaker tests (echoic and tact) were structured similarly to the pre-tests previously described. The blocks consisted of three trials for each stimulus in echoic (in a total of nine echoic trials) and for three more tact trials (a total amount of nine tact trials), randomized into a single block in a total of 18 trials. If picture naming behavior did not emerge following the speaker tests with stimuli from Set 1, the participants were exposed to the teaching through MEI with stimuli from Set 2.

MEI – Multiple Exemplar Instruction (Set 2). The MEI consisted of teaching listening (pointing) and speaker (echoic and tact) skills, in rotation, using Set 2 stimuli. During MEI, programmed consequences were provided for correct or incorrect responses. Error correction for listening tasks were identical to those described in listening training with Set 1, with a gestural prompt; for echoic and tact behaviors, an echoic prompt was provided, that is, the word was repeated by the experimenter, along with an orofacial cue, which allowed the child to observe the researcher’s vocal topography; then, another opportunity to respond was given. MEI training was conducted in a linear manner, that is, the three types of trials were presented for each stimulus in Set 2, in the sequence: echoic, pointing to, and tact. This instructional sequence is presented in Olaff e Holth (2025) as Mixed-Operant Instruction (MOI). Teaching began with an echoic trial in which an auditory stimulus (i.e., “ball”) was presented, and the participant was encouraged to vocalize the word.

After repeating the word and applying correction procedures, if necessary, three pictures were presented to the participant, one of which corresponded to “ball”, for example, then, the experimenter said, “Point to... (ball)”. After selecting the picture, the consequences for correct and incorrect responses were provided, and correction procedures were applied. During the tact phase, the ball picture was presented, and the participant emitted the tact. After the programmed consequence for correct and incorrect responses was provided, and the correction procedure for the previous trial (involving a ball) was carried out, the experimenter presented another stimulus, following the same procedures. Each block consisted of 27 trials, nine for listening and speaker (echoic and tact) behaviors, and three for each stimulus in the set. The echoic, listener, and tact trials targeted the same items, and the items were randomized into the blocks of trials.

Listening training and Speaker tests review (Set 1). After MEI, the listener training was reviewed, and probes of the speaker responses with stimuli from Set 1 were conducted again.

Listening training (pointing to) and speaker test (echoic and tact) (Set 3). The same structure applied to the listener training in Set 1.

Post-test. This phase was procedurally identical to the Pre-test.

Sessions Routine

A routine that comprises teaching and testing procedures and reinforcing conditions was applied to the sessions. An individual routine was adopted for each participant, according to their initial repertoire. The information was collected from interviews with those responsible for obtaining possible reinforcing items to be used in the sessions (for John, his mother; for Harry, his grandmother).

The routine started with an oral instruction regarding the execution of the routine steps, such as “Let’s have two types of activities: one on the computer, with sounds and pictures, and the other one is a game. Let’s take turns on the two types of activities”. Two researchers were present while the routine was carried out. One of them presented the trials during teaching and testing phases via laptop, and the other one recorded the participants’ responses and organized the reinforcing activities that came along with the corresponding phase of the procedure. The routine ended by concluding the activities with a game, putting the toys away, and saying “*Goodbye*”. The activities were organized as a circuit in the activity room. After each step of the program, there was a pause to play a game for about 10 min. The tasks of the day were also concluded with a game. Data collection took place for a week, with two sessions per day, lasting 50 min each.

Procedure for results analysis

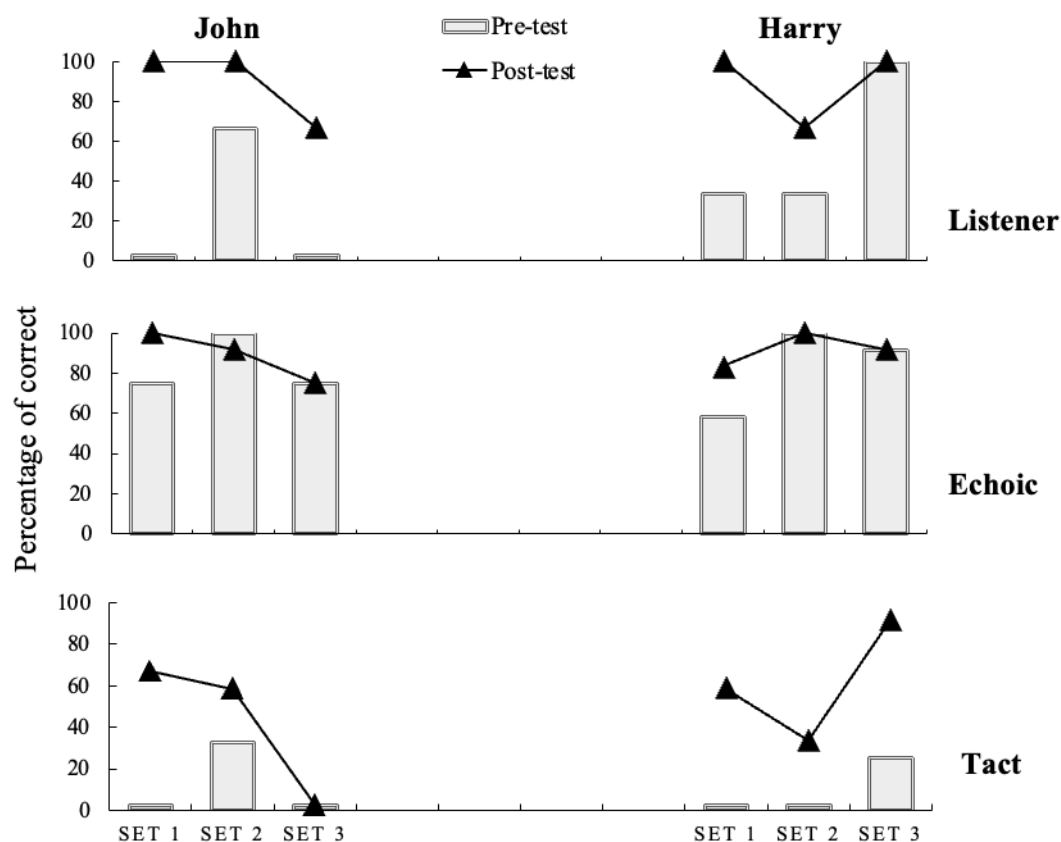
Listener behavior was analyzed as the percentage of pictures selected correctly and independently during each trial block. Speaker behavior was assessed through the transcription of the vocal responses that were emitted by the participants and their analysis in accordance with the percentage of correct responses and their point-by-point correspondence to normative speech; for that, a phonemic analysis was carried out (Barreto & Ortiz, 2008). The vocal responses (picture tact and echoic behaviors) were analyzed by two researchers. At the end of each session, the researchers reached an agreement regarding the responses that were transcribed and registered. The vocalizations were scored following the point-to-point correspondence to verbal community conventions. Researchers obtained a rate that related the phonemes emitted correctly and the phonemes defined as correct. For example, if the objective was to emit “bola”, and the participant vocalized /nola/, it corresponded to three correct phonograms /ola/ and an incorrect one /n/; therefore, such response was scored with 75% of accuracy; if the participant vocalized /po--/, was scored 25% of accuracy; if no was vocalized, it was scored as zero. Thus, the emergence of the speaker's behavior was analyzed separately from accuracy (correct vocal topography), which was measured using the phonemic score.

Results

Figure 2 demonstrates the participants’ performances in the Pre-test (bars) and Post-test assessments of listener, echoic, and tact repertoires. The best results were in echoic behavior: from 70% to 100% for John, and from 60% to 100% for Harry. The lowest results were with tacting: from zero to 40% for John (best result in Set 2), and from zero to 30% for Harry (best result in Set 3). For listener behavior, the results were from zero to 60% for John (best result in Set 2), and Harry demonstrated variability with intermediate values, from 30% to 100% (best result in Set 3). Overall, the participants’ percentage of correct responses across listener, echoic, and tact repertoires improved Post-MEI training. Especially for listener and tact behaviors, responding was zero or close to zero before MEI training and increased after it. The participants achieved 100% correct responses in listener behavior in two sets of stimuli. With echoic behavior, correct response percentages ranged between 80% and 100% accuracy. With tacting, the percentages were above the baseline levels for both participants, except for John in Set 3 (Figure 2). The data for correct listener, echoic, and tact behaviors for John and Harry, are displayed in Figure 3. In the Pre-test phase, tact behavior was absent (John) or very low, below 25% (Harry) for stimuli of Sets 1 and 3. Before MEI, listener training was not enough for a significant improvement (Set 1). The percentage of correct responses for tact behavior increased only after MEI (Sets 1 and 3; Figure 3)

Figure 2

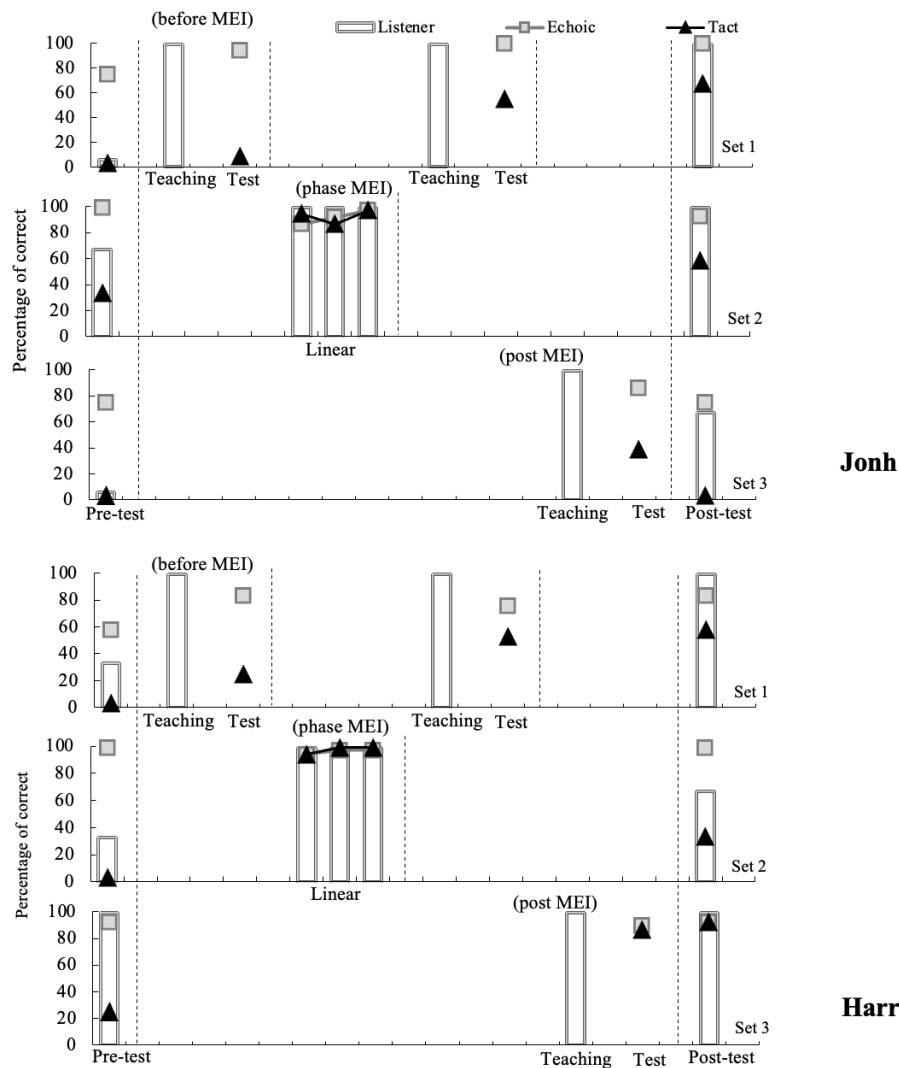
General Percentage of Correct in Listener, Echoic and Tact Repertoires During Pre and Post-tests



Note. Bars represent Pre-test performance, and triangles connected by a line represent posttest performance.

Specifically, according to Figure 3, in the Pre-test phase, John did not engage in listener responses for Sets 1 and 3, and accurate responses were at 67% for Set 2. For speaker responses, John's correct responding was at 33% for Set 2; echoic responses were at 75% for Sets 1 and 3; and 100% in Set 2. These results demonstrate the variability of dependent variables across stimuli sets, and the independence between listener and speaker behaviors. John's listener responses were at 100% after the first exposure to selection-based listening training with Set 1. For his speaker responses, echoic behavior was at 94% and tact, at 8%. Listener training with Set 1, by itself, was not enough to produce emergent speaker responses (tact and echoic).

Three exposures to the MEI procedure with Set 2 were necessary for John to emit echoic and tact behaviors with 100% accuracy. Speaker responses were assessed after listener training with Set 1, to ensure responding was still at 100%. Echoic behavior accuracy was 100% and tact behavior was at 55%. In the listener training with a new set of stimuli (Set 3), only one exposure was necessary for John to achieve 100% accuracy. After this training, echoic accuracy was 86% and tact was at 39%, an increase above Pre-MEI probe levels. Responding was less variable in Post-MEI probes compared to Pre-MEI probes, from 65% to 100% for echoic and listener behavior. For tact, the results improved Post-MEI training compared to Pre-MEI, for Sets 1 and 2, with 67% and 58% accuracy, respectively.

Figure 3. *Percentage of Correct Responses in Listener, Echoic, and Tact Repertoires Before and After MEI*

Note. Data are shown for John (top) and Harry (bottom) across Sets 1, 2, and 3. Dashed vertical lines indicate phase changes. Listener responses are shown as bars, echoic responses as squares, and tact responses as triangles.

The results for Harry demonstrated that the listener response for Sets 1 and 2 was at 33%, and for Set 3, 100%. He did not present tact behavior for Sets 1 and 2, and, for Set 3, his tact behavior was at 25%. As for echoic behavior, the responses were at 58% for Set 1, 100% for Set 2, and 92% for Set 3. For Set 1, Harry's listener response was 100% accurate after the first exposure to listener training. After listener training, his echoic behavior improved in comparison to pre-listener training levels at 80%; the same happened to tacting, which was 25% accurate after the training. After starting MEI with Harry, listener responses were 100% accurate during three exposures to MEI. The speaker results were like echoic responding, at 94%, 97%, and 97%; and tacting at 94%; 100%, and 100%. After MEI, listener responding was tested again for Set 1, to ensure responding was still at 100%. During the tests for speaker repertoires with stimuli from the same set (Set 1) and a new set (Set 3), there was an increase in the percentage of correct responses for echoic and tact behaviors. Therefore, the percentages of correct responses were higher than those in the pre-tests phase, getting to 75% correct responses in echoic behavior and 53% for tacting.

Discussion

This study demonstrated the effectiveness of the MEI procedure to increase the percentage of correct responses in speaker and listening responses (Figure 3, Set 2, in the middle) in two children diagnosed with Auditory Neuropathy Spectrum Disorder (ANSD). This procedure was successful in promoting emergent speaker responding

after listener training (Figure 3, Set 1 and Set 3, upper and lower right). Before exposure to MEI, selection-based listening behavior was not sufficient to produce tact emergence. After exposure to MEI, an increase in tacting was observed compared to performance prior to listener training, which is consistent with greater functional integration between listener and speaker repertoires after MEI.

The results of this study align with the results of a previous study, conducted with one participant diagnosed with ANSD (Mascotti et al. 2024), in which the functional interdependence between the listening and speaker responses was achieved during the MEI phase. Following MEI, the participants received listening training with a new set of stimuli, and the results in speaker (tact and echoic) behavior improved. In the present study, the MEI procedure was implemented linearly for each stimulus (cf. Olaff & Holth, 2025), rotating across verbal operants and listener responding within the same teaching sequence. In other words, for the same item, the participant was successively exposed to echoic responses ("Repeat..."), listener responses ("Point to..."), and tacting "The name is..."). This instructional sequence was relevant because, in addition to involving multiple operants, it promoted shared control between stimuli (the stimulus that controlled the echoic response was the same one that controlled the listener response) and between responses (the topography emitted in the echoic response was the same as that emitted in the tactile response), favoring the interdependence between hearing and speaking. In the Pre- and Post-tests, the operant and items were presented in randomized supporting generalization measure (Olaff & Holth, 2025).

The MEI procedure resulted in emergent speaker behavior in both participants. The emergence of tacting after the listener training (Figure 3, at the bottom) indicates that MEI (which included listener, echoic, and tact responses) contributes to speaker unidirectional naming. In the present study, this result is described as unidirectional speaker naming, that is, tact emergence after listener instruction and after exposure to the MEI. This does not correspond to a demonstration of bidirectional naming, which would require additional evidence of tact and listening emergence after exposure to an indirect and incidental experience of these repertoires (Hawkins et al., 2018; Sivaraman & Barnes-Holmes, 2023), representing an important gap for conducting future studies with this population.

The alternation between echoic and tact trials contributes to stimulus control over operants of the same vocal topography. In other words, while echoic training promotes improvements in the accuracy of emission of the vocalized word, it also promotes responding under picture control. In addition, selection-based listener training strengthens the relation between dictated words and pictures, enabling picture naming (Horne et al., 2006; Santos & Souza, 2016).

This interpretation is also strengthened by recent work on the role of echoic behavior. Frampton and Miguel (2025) argue that when transfer of control, from listening to speaker repertoire is observed, this likely occurs, in part, because the child echoes the instruction during listener behavior, and reinforcement contingent on the selection response also strengthens vocalization in the presence of the nonverbal stimulus. Similarly, Engell et al. (2025) indicate that echoic behavior may have an important facilitative role in the emergence of tacting, although the results are not yet uniform across experimental evaluations. Therefore, in the present study, echoic behavior can be treated as a relevant variable for the transfer of control, without this implying attributing to it, in isolation, the entire observed effect.

Although there is variability in outcomes between the participants, the results after MEI increased during post-tests phase, above the baseline levels, for the three target behaviors (Figure 2), and can be said that MEI reduced the differences between listening and speaking, which demonstrated more integration during Post-tests (Figure 3).

Similar results to the current study were also reported by Silva et al. (2024). In their study, after the implementation of the MEI procedure involving listener and tact training, a typically developing and hearing child (aged 28-29 months) demonstrated unidirectional speaker naming. The other two children demonstrated improved performances (between 40% to 60% accuracy) compared to the baseline phase (in which there were no responses). In contrast to this study, Silva et al. (2024) did not include echoic training in the MEI procedure. In the current study, the lack of accuracy in the vocalizations observed in emergent tact responses (assessed in Set 3) may be partly related to articulation difficulties; however, during MEI (Set 2), when the echoic model was available, tact responses were more accurate. After MEI, when both echoic and tact performances were tested, vocalizations became less accurate, although the percentage of correct responses remained higher than in the Pre-tests phase. These results suggest a dependency on auditory stimulus control for accurate vocalizations to be issued. Furthermore, the percentage of correct answers in speaker behavior was obtained according to speech accuracy. This means that the participant may have vocalized "ball" instead of "paw" or "cat" instead of "bat"; since there were no pictures of ball or cat as stimuli, these vocalizations can be considered partially correct, as they closely approximate the target vocalization and were not emitted at baseline.

In this sense, the procedure can be described as efficient because it produced more accurate speech. However, when considering speech accuracy, this analysis reduced the percentage of correct answers, since, except for Harry in

Set 3, none of the participants obtained 100% or near-100% accuracy in the speaker measures. The fact that the participants had high percentages of correct answers in echoic responses for stimuli from all three sets from baseline, and that they improved throughout the stages of the procedure, ensures that they could emit the phonemes (Souza & Calandrini, 2022) with better accuracy than with tacting, even though the required topography was the same. This discussion should be interpreted in conjunction with the methodological limitation regarding the absence of permanent audiovisual recordings. Since the dependent variable included the correct vocal topography, the lack of stable recordings reduces the robustness of the point-to-point correspondence classification between the auditory model and the emitted response. Conditions for stimulus control transfer from echoic to tact should be further investigated in future studies with this population. Another limitation is the absence of repeated baseline and follow-up measurements; repeated baseline probes did not allow verification of the stability of vocalizations at very low levels before the intervention; the absence of follow-up did not allow verification of the maintenance of the repertoire over time and in the absence of experimental conditions. Future research should plan for repeated monitoring of tact emergence, even in clinical circumstances such as those in the present study, given that the arrangement of these conditions represents a gap.

The two participants presented considerable differences in age and in time of Cochlear Implant (CI) use, which suggests differences regarding the exposure to auditory stimulation, and, therefore, in the development of speaker verbal behavior (Geers et al., 2011; Geers et al., 2016). Despite such differences, John's and Harry's performances during the training phase were similar and indicate that MEI can be applied with individuals who demonstrate limited verbal behavior and with variability, and able to reduce differences in acquisition within and across individual learners. Understanding the conditions under which the variability between listening and speaker repertoires is minimized is fundamental for building a teaching technology that produces relevant behavior changes (Conceição et al., 2022). In future studies, it may be helpful to report and evaluate information about the pre-implanted Pure-Tone Audiometry (PTA) or the type of cochlear implant, to check how the MEI procedure can be used to address individual performances with even greater precision.

The results for speaker behavior with stimuli from Set 3 were particularly interesting (Figure 3, bottom, on the right). Such results are different from the ones obtained with Set 1 because these stimuli were trained before and after the MEI, which could affect the results, as there was more exposure to training. A previous study probed selection behavior (i.e., listening to a word and pointing at pictures) that was followed by an assessment of labeling, and, at first, the participants did not demonstrate this skill. However, repeated probes regarding the speaker behavior, by itself, resulted in increases in labeling behavior (Bandini et al., 2012). Nevertheless, the results of stimuli from Set 3 in this study might represent a measure of generalization regarding the effects of MEI for untrained sets of stimuli. If echoic is well established, listening to a new word might generate subvocal echoic responses. Moreover, listening to a new word and associating it with a picture, in a circumstance in which only the picture is presented, might lead to tacting responses later (Geers et al., 2016; Greer & Ross, 2008; Santos & Souza, 2020). Including this third set in stages that followed MEI was critical and presented only in a previous study (Mascotti, 2019); also, it had not yet been measured in other previous research that applied MEI for children with hearing disabilities (Merlin, et al., 2019; Pereira et al., 2016; Rique et al., 2017).

The role of Set 3 in the present study was to function as a control condition and to verify whether, after exposure to the MEI, a new listener teaching sequence, conducted after the integrated structure between listening and speaking by the MEI, would favor the emergence of speaker responses. As the results were positive, the possibility of interpreting the speaker results with stimuli from Set 1 because of overtraining is reduced. Data collection was conducted in a format that is typical for common clinical conditions, which represents a positive feature of this study. It suggests that the procedures evaluated in this study are suitable to typical clinical sessions and, therefore, are more likely to be adopted by clinicians (Grow & Kodak, 2010). Moreover, the variety of environments in which the research studies were carried out might reveal an ecological dimension to the study if one considers that the more the conditions for data collection get close to the everyday life of the participants, the greater the potential for ecological validity of the experimental situation and external validity of the results.

MEI is a teaching procedure that more closely approximates the contingencies present in a child's daily life than procedures that rigidly isolate the verbal operant. In everyday language use, listening, echoing, and tacting do not usually occur as independent repertoires, but in a functional and often integrated relationship. In clinical protocols, this suggests that the MEI can be used when isolated listener training is not, by itself, producing a transfer of control to vocal responses. The study by Olaff and Holth (2025), for example, showed that the type of probe, the order of the responses tested, and the repetition of probes can alter the observed pattern of emergence and maintenance, which reinforces

the importance of adjusting these parameters in future research and clinical protocols. MEI procedure is flexible, and, beyond vocal repertoires, it also leads to the acquisition of new motor and sensorial repertoires. Such benefits have been discussed in previous studies that assessed several target responses, such as dictation and verbal interactions. Some of them include: answering questions (Conceição et al., 2022; Greer et al., 2005); word repetition; picture naming; and asking for things, in children with ASD (Guerra & Almeida Verdu, 2020); emission of sentences with adjectives, in typically developing children (Luke et al., 2011) and with cochlear implant (Merlin et al., 2019); and listening responses and verbal interactions, in children with ASD (Lechago et al., 2015).

As for the study limitations, three aspects are noteworthy. One of them refers to the lack of videos and stable records for further analysis of the vocalizations, as well as the absence of follow-up data collection to evaluate the maintenance of speaker and listening responses. This emergent behavior is called incidental bidirectional naming (Hawkins et al., 2018; Santos & Souza, 2016) and has been demonstrated in children with hearing loss and cochlear implants only once, in a study with one participant (Pereira et al. 2018). Currently, there is an open conceptual and methodological debate in the bidirectional naming field (Krüger et al. 2024). This study intended to replicate previous studies with children with hearing impairment in a clinical context, however, in participants with ANSD and CI, the evidence is still preliminary, and the findings of this study should be understood as a promising expansion of this line of investigation. Recent literature emphasizes that most studies on MEI and naming still focus on children with typical development or ASD, and that extending these results to populations with hearing loss and cochlear implants still requires further research.

Future research should explore (cf. Olaff & Holth, 2025; Santos & Souza, 2020; Sivaraman & Barnes-Holmes, 2023): (a) the specific contribution of echoic to the topographic accuracy of tact; (b) the monitoring of tact emergence including repeated probes throughout the study, including audiovisual recording for subsequent inter-observer agreement (IOA) analysis; (c) the effect of probe order and density on observed emergence; (d) the maintenance of the repertoire after non-training intervals; and (e) the extension of the procedure to populations with ANSD using CI in larger samples.

Conflict of Interest Statement

The authors declare that there are no conflicts of interest regarding the publication of this article.

Contribuição de cada autor

We certify that all authors have participated sufficiently in the work to take public responsibility for its content. Each author's contribution can be attributed as follows: Ana Claudia M. Almeida-Verdu: definition of the research problem, methodological strategy, organization and interpretation of results, review of the manuscript. Aline Daniela G. S. Vieira: Interpretation of results, writing of the manuscript, formatting in accordance with APA standards, and review of bibliographic references. Fernanda de Oliveira Ortigoza: Interpretation of results, writing of the manuscript. Felipe Monteiro-Silva: Interpretation of results, writing of the manuscript and review of bibliographic references. Maria Clara L. Briseno: Interpretation of results, writing of the manuscript. Sarah A. Lechago: Discussion of data, review of the manuscript, and adaptation of the academic language in the area in English.

Copyright

This is an open-access article and can be freely reproduced, distributed, transmitted, or modified by anyone, provided it is used for non-commercial purposes. The work is made available under the Creative Commons 4.0 BY-NC license.



References

- Almeida Verdu, A. C. M., da Silva, W. R., Golfeto, R. M., Bevilacqua, M. C., & de Souza, D. G. (2014). Investigação da função simbólica adquirida por estímulos elétricos em crianças com implante coclear. In: *Comportamento simbólico: bases conceituais e empíricas*. Marília: Oficina Universitária. pp. 229-268 <https://doi.org/10.36311/2014.978-85-7983-516-2.p229-267>
- Bandini, C. S. M., Sella, A. C., Postalli, L. M. M., Bandini, H. H. M., & Silva, E. T. P. (2012). Effects of selection tasks on naming emergence in children. *Psicologia: Reflexão e Crítica*, 25(3), 568 – 577. <https://doi.org/10.1590/S0102-79722012000300017>

- Barreto, S. dos S., & Ortiz, K. Z. (2008). Medidas de inteligibilidade nos distúrbios da fala: revisão crítica da literatura. *Pró-Fono Revista de Atualização Científica*, 20(3), 201–6. <https://doi.org/10.1590/S0104-56872008000300011>
- Carvalho, A. C. M., Bevilacqua, M. C., Samechima, K., & Costa, A. O. (2011). Auditory neuropathy/Auditory dyssynchrony in children with cochlear implants. *Brazilian Journal of Otorhinolaryngology*, 77, 481-7. <https://doi.org/10.1590/S1808-86942011000400012>
- Conceição, D., B. da, Greer, R. D., & Lee-Moschella, J. (2022). A general outline of the Verbal Behavior Developmental Theory. *Revista Brasileira de Terapia Comportamental e Cognitiva*, 24, 1-39. <https://doi.org/10.31505/rbtcc.v24i1.1646>
- Costa, N. T. O., Martinho-Carvalho, A. C., Cunha, M. C., & Lewis, D. R. (2012). Habilidades auditivas e comunicativas no espectro da neuropatia auditiva e mutação no gene otoferlin: Estudo de casos. *Jornal da Sociedade Brasileira de Fonoaudiologia*, 24(2), 181-187. <https://doi.org/10.1590/S2179-64912012000200016>
- Engell, T. S. (2024). *The role of the echoic in the acquisition of tacts* [California State University, Sacramento]. <https://scholars.csus.edu/esploro/outputs/graduate/The-role-of-the-echoic-in/99258173463201671#file-0>
- Fernandes, N. F., Morettin, M., Yamaguti, E. H., Costa, O. A., & Bevilacqua, M. C. (2015). Performance of hearing skills in children with auditory neuropathy spectrum disorder using cochlear implant: a systematic review. *Brazilian Journal of Otorhinolaryngology*, 81(1), 85–96. <https://doi.org/10.1016/j.bjorl.2014.10.003>
- Frampton, S. E., & Miguel, C. F. (2025). [Generative and emergent verbal behavior](#). In: Vladescu, J.C., & Kisamore, A.N. (Eds.). *Promoting language for learners with autism spectrum disorder: A verbal behavior guide for practitioners* (1st ed.). Routledge. <https://doi.org/10.4324/9781003433668>
- Geers, A. E., & Hayes, H. (2011). Reading, writing, and phonological processing skills of adolescents with 10 or more years of cochlear implant experience. *Ear and Hearing*, 32(1), 49S-59S. <https://doi.org/10.1097/AUD.0b013e3181fa41fa>
- Geers, A. E., Nicholas, J., Tobey, E., & Davidson, L. (2016). Persistent language delay versus late language emergence in children with early cochlear implantation. *Journal of Speech, Language, and Hearing Research*, 59(1), 155–170. https://doi.org/10.1044/2015_JSLHR-H-14-0173
- Gilic, L., & Greer, R. D. (2011). Establishing naming in typically developing two-year-old children as a function of multiple exemplar speaker and listener experiences. *The Analysis of Verbal Behavior*, 27, 157-177. <https://doi.org/10.1007/BF03393099>
- Greer, R. D., Stolfi, L., Chavez Brown, M., & Rivera Valdes, C. (2005). The emergence of the listener to speaker component of naming in children as a function of multiple exemplar instruction. *The Analysis of Verbal Behavior*, 21(1), 123 - 134. <https://doi.org/10.1007/BF03393014>
- Greer, R. D., & Ross, D. E. (2008). *Verbal Behavior Analysis: Inducing and expanding new verbal capabilities in children with language delays*. Boston: Pearson.
- Greer, R. D., & Speckman, J. (2009). The integration of speaker and listener responses: A theory of verbal development. *The Psychological Record*, 54, 449-488. <https://doi.org/10.1007/BF03395674>
- Grow, L. L., & Kodak, T. (2010). Recent research on emergent verbal behavior: Clinical applications and future directions. *Journal of Applied Behavior Analysis*, 43(4), 775-778. <https://doi.org/10.1901/jaba.2010.43-775>
- Guerra, B. T., & Almeida-Verdu, A. C. M. (2020). Ensino de comportamento verbal elementar por exemplares múltiplos em crianças com autismo. *Psicologia: Ciência e Profissão*, 40, 1-17. <https://doi.org/10.1590/1982-3703003185295>
- Guess, D. (1969). A functional analysis of receptive language and productive speech: acquisition of the plural morpheme. *Journal of Applied Behavior Analysis*, 2, 55-64. <https://doi.org/10.1901/jaba.1969.2-55>
- Hawkins, E., Kingsdorf, S., Charnock, J., Szabo, M., & Gautreaux, G. (2009). Effects of multiple exemplar instruction on naming. *European Journal of Behavior Analysis*, 10(2), 265-273. <https://doi.org/10.1080/15021149.2009.11434324>
- Hawkins, E., Gautreaux, G., & Chiesa, M. (2018). Deconstructing common bidirectional naming: a proposed classification framework. *The Analysis of Verbal Behavior*, 17(34), 44-61. <https://doi.org/10.1901/jeab.1996.65-185>
- Horne, P. J., & Lowe, C. F. (1996). On the origins of naming and other symbolic behavior. *Journal of the Experimental Analysis of Behavior*, 65(1), 185–241. <https://doi.org/10.1901/jeab.1996.65-185>
- Horne, P., Hughes, C., & Lowe, F. (2006). Naming and categorization in young children: IV: Listener behavior training and transfer of function. *Journal of the Experimental Analysis of Behavior*, 85(2), 247-273. <https://doi.org/10.1901/jeab.2006.125-04>
- Hussein, L. G., Góes, C. C., Chiodelli, T., Silva-Marinho, C. S. O., Gonçalves, F. L., & Almeida Verdu, A. C. M. (2018). Aquisição do comportamento de ouvir, baseada em seleção de figuras, em crianças com implante coclear contralateral. *Revista Brasileira de Terapia Comportamental e Cognitiva*, 10(1), 27-39. <https://doi.org/10.31505/rbtcc.v20i1.1135>

- Krüger, G., da Conceição, D., & de Rose, J. (2024). Uma introdução à Nomeação Bidirecional. *Revista Brasileira de Análise do Comportamento*, 20. <https://doi.org/10.18542/rebac.v20i0.16458>
- Lechago, S. A., Carr, J. E., Kisamore, A. N., & Grow, L. L. (2015). The effects of multiple exemplar instruction on the relation between listener and intraverbal categorization repertoires. *The Analysis of Verbal Behavior*, 31(1), 76-95. <https://doi.org/10.1007/s40616-015-0027-1>
- Longano, J., & Greer, R. D. (2015). Is the source of reinforcement for naming multiple conditioned reinforcers for observing responses? *The Analysis of Verbal Behavior*, 31(1), 96-117. <https://doi.org/10.1007/s40616-014-0022-y>
- Luke, N., Greer, R. D., Singer Dudek, J., & Keohane, D. D. (2011). The emergence of autoclitic frames in atypically and typically developing children as a function of multiple exemplar instruction. *The Analysis of Verbal Behavior*, 27(1), 141-156. <https://doi.org/10.1007/BF03393098>
- Mascotti, T. de S., Almeida-Verdu, A. C. M., Silva, L. T. N., LaFrance, D., & McIlvane, W. J. (2024). Emergência de tato e precisão da fala via Instrução por Múltiplos Exemplos em crianças com pouco repertório verbal. *Acta Comportamental*, 32(3), 423-447. <https://doi.org/10.32870/ac.v32i3.88365>
- Mascotti, T. S., & Almeida Verdu, A. C. M. (2020). Lapidar: Programa de ensino, ampliação e refinamento de comportamento verbal. Bauru, SP: Faculdade de Ciências, Universidade Estadual Paulista. <https://sgcd.fc.unesp.br/#!/lads/producao/produto-tecnico-e-tecnologico-programas/lapidar/>
- Merlin, A. M. B., Almeida-Verdu, A. C. M., Neves, A. J. da, Silva, L. T. do N., & Moret, A. L. M. (2019). Ensino por múltiplos exemplares e integração de comportamentos de ouvinte e falante com unidades sintáticas substantivo-adjetivo em crianças com DENA e IC. *CoDAS*. 31(3):e20180135. <https://doi.org/10.1590/2317-1782/20182018135>
- Miguel, C., Petursdottir, A. I., Carr, J. E., & Michael, J. (2008). The role of naming in stimulus categorization by preschool children. *Journal of Experimental Analysis of Behavior*, 89(3), 383-405. <https://doi.org/10.1901/jeab.2008-89-383>
- Morin, K. L., Lindström, E. R., Kratochwill, T. R., Levin, J. R., Blasko, A., Weir, A., Nielsen-Pheiffer, C. M., Kelly, S., Janunts, D., & Hong, E. R. (2024). Nonconcurrent multiple-baseline and multiple-probe designs in special education: a systematic review of current practice and future directions. *Exceptional Children*, 90(2), 126-147. <https://doi.org/10.31219/osf.io/r6zsa>
- Norrix, L. W., & Velenovsky, D. S. (2014). Auditory neuropathy spectrum disorder: a review. *Journal of Speech, Language and Hearing Research*, 57(4), 1564-76. https://doi.org/10.1044/2014_JSLHR-H-13-0213
- Olaff, H. S., Ona, H. N., & Holth, P. (2017). Establishment of naming in children with autism through multiple response-exemplar training. *Behavioral Development Bulletin*, 22(1), 67-85. <https://doi.org/10.1037/bdb0000044>
- Olaff, H. S., & Holth, P. (2025). Acquisition of incidental bidirectional naming: Isolating the effects of probing and mixed-operant instruction. *The Analysis of Verbal Behavior*, 41, 200-234. <https://doi.org/10.1007/s40616-025-00221-1>
- Pereira, F. D., Assis, G. J. A., & Almeida Verdu, A. C. M. (2016). Integração dos repertórios de falante-ouvinte via instrução com exemplares múltiplos em crianças implantadas cocleares. *Revista Brasileira de Análise do Comportamento*, 12(1), 23-32. <https://doi.org/10.18542/rebac.v12i1.4023>
- Pereira, F. S., Assis, G. J. A., Neto, F. X. P., & Almeida Verdu, A. C. M. (2018). Emergência de nomeação bidirecional em criança com implante coclear via Instrução com Múltiplos Exemplos (MEI). *Revista Brasileira de Terapia Comportamental e Cognitiva*, 20(2), 26-39. <https://doi.org/10.31505/rbtcc.v20i2.1178>
- Rique, L. D., Guerra, B. T., Borelli, L. M., Oliveira, A. P., & Almeida Verdu, A. C. M. (2017). Ensino de comportamento verbal por múltiplos exemplares em uma criança com distúrbio do espectro da neuropatia auditiva. *CEFAC*, 19(2), 289-298. <https://doi.org/10.1590/1982-021620171928516>
- Santos, E. L. N., & Souza, C. B. A. (2016). Ensino de nomeação com objetos e figuras para crianças com autismo. *Psicologia: Teoria e Pesquisa*, 32(3), 1-10. <https://doi.org/10.1590/0102-3772e32329>
- Santos, E. L. N., & Souza, C. B. A. (2020). Uma Revisão Sistemática de Estudos Experimentais sobre Nomeação Bidirecional. *Revista Brasileira de Análise do Comportamento*, 16(2). <https://doi.org/10.18542/rebac.v16i2.9605>
- Silva, G. G. da, Souza, C. B. A. de, & Gil, M. S. C. de A. (2024). Unidirectional speaker naming in toddlers: effects of Multiple Exemplar Instruction teaching. *Paidéia (Ribeirão Preto)*, 34, e3415. <https://doi.org/10.1590/1982-4327e3415>
- Sivaraman, M., & Barnes-Holmes, D. (2023). Naming: What do we know so far? A systematic review. *Perspectives on Behavior Science*, 46(3-4), 585-615. <https://doi.org/10.1007/s40614-023-00374-1>
- Skinner, B. F. (1957/1978). *O comportamento verbal*. São Paulo: Cultrix.
- Souza, C. B. A., & Calandrini, L. (2022). Pareamento de estímulos e aquisição de comportamento verbal em crianças com TEA. *Acta Comportamental*, 30(1). <https://doi.org/10.32870/ac.v30i1.81397>

Yoder, P., Camarata, S., & Gardner, E. (2005). Treatment Effects on Speech Intelligibility and Length of Utterance in Children with Specific Language and Intelligibility Impairments. *Journal of Early Intervention*, 28(1), 34-49. <https://doi.org/10.1177/105381510502800>

Submitted: 10/07/2025

Accepted: 07/05/2026