



Brazilian Journal of OTORHINOLARYNGOLOGY

www.bjorl.org



ARTIGO ORIGINAL

Encoding of speech sounds at auditory brainstem level in good and poor hearing aid performers[☆]



Hemanth Narayan Shetty* e Manjula Puttabasappa

All India Institute of Speech and Hearing, Department of Audiology, Mysuru, Karnataka, Índia

Recebido em 5 de fevereiro de 2016; aceito em 20 de junho de 2016

Disponível na Internet em 29 de junho de 2017

KEYWORDS

Frequency following
response;
Acceptable noise
level;
Hearing aid
performer

Abstract

Introduction: Hearing aids are prescribed to alleviate loss of audibility. It has been reported that about 31% of hearing aid users reject their own hearing aid because of annoyance towards background noise. The reason for dissatisfaction can be located anywhere from the hearing aid microphone till the integrity of neurons along the auditory pathway.

Objectives: To measure spectra from the output of hearing aid at the ear canal level and frequency following response recorded at the auditory brainstem from individuals with hearing impairment.

Methods: A total of sixty participants having moderate sensorineural hearing impairment with age range from 15 to 65 years were involved. Each participant was classified as either Good or Poor Hearing aid Performers based on acceptable noise level measure. Stimuli/da/and/si/were presented through loudspeaker at 65 dB SPL. At the ear canal, the spectra were measured in the unaided and aided conditions. At auditory brainstem, frequency following response were recorded to the same stimuli from the participants.

Results: Spectrum measured in each condition at ear canal was same in good hearing aid performers and poor hearing aid performers. At brainstem level, better F_0 encoding; F_0 and F_1 energies were significantly higher in good hearing aid performers than in poor hearing aid performers. Though the hearing aid spectra were almost same between good hearing aid performers and poor hearing aid performers, subtle physiological variations exist at the auditory brainstem.

Conclusion: The result of the present study suggests that neural encoding of speech sound at the brainstem level might be mediated distinctly in good hearing aid performers from that of poor hearing aid performers. Thus, it can be inferred that subtle physiological changes are evident at the auditory brainstem in a person who is willing to accept noise from those who are not willing to accept noise.

© 2016 Associação Brasileira de Otorrinolaringologia e Cirurgia Cérvico-Facial. Published by Elsevier Editora Ltda. This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

DOI se refere ao artigo: <http://dx.doi.org/10.1016/j.bjorl.2016.06.004>

[☆] Como citar este artigo: Shetty HN, Puttabasappa M. Encoding of speech sounds at auditory brainstem level in good and poor hearing aid performers. Braz J Otorhinolaryngol. 2017;83:512–22.

* Autor para correspondência.

E-mail: hemanthn.shetty@gmail.com (H.N. Shetty).

A revisão por pares é da responsabilidade da Associação Brasileira de Otorrinolaringologia e Cirurgia Cérvico-Facial.

2530-0539/© 2016 Associação Brasileira de Otorrinolaringologia e Cirurgia Cérvico-Facial. Publicado por Elsevier Editora Ltda. Este é um artigo Open Access sob uma licença CC BY (<http://creativecommons.org/licenses/by/4.0/>).

PALAVRAS-CHAVE

Frequência seguida de resposta;
Nível de ruído aceitável;
Usuário de aparelho auditivo

Codificação dos sons da fala no nível do tronco encefálico auditivo em bons e maus usuários de aparelhos auditivos**Resumo**

Introdução: Os aparelhos auditivos são prescritos para aliviar a perda de audibilidade. Tem sido relatado que 31% dos usuários rejeitam seu aparelho auditivo devido ao desconforto com o ruído de fundo. A razão para a insatisfação pode estar situada em qualquer local desde o microfone do aparelho auditivo até a integridade de neurônios ao longo da via auditiva.

Objetivos: Medir espectros desde a saída do aparelho auditivo no nível do meato acústico externo e frequência de resposta (FFR) registrada no tronco encefálico de indivíduos com deficiência auditiva.

Método: Foram selecionados 60 participantes com deficiência auditiva neurossensorial moderada, de 15 a 65 anos. Cada participante foi classificado como usuário bom ou mau de prótese auditiva (GHP ou PHP) com base na medida de nível de ruído aceitável (ANL). Estímulos /da/e/si/ foram apresentados em alto-falante a 65 dB SPL. No meato acústico externo, os espectros foram medidos nas condições sem aparelho e com aparelho. No tronco encefálico auditivo, FFR foram registradas para os mesmos estímulos dos participantes.

Resultados: Os espectros medidos em cada condição no meato acústico externo foram os mesmos em GHP e PHP. No nível do tronco cerebral, melhor codificação F0; energias de F0 e F1 foram significativamente maiores em GHP do que em PHP. Embora os espectros do aparelho auditivo fossem quase os mesmos entre GHP e PHP, existem variações fisiológicas sutis no tronco encefálico auditivo.

Conclusão: O resultado do presente estudo sugere que a codificação neural do som da fala no nível do tronco encefálico pode ser mediada distintamente em GHP em comparação com PHP. Assim, pode-se inferir que mudanças fisiológicas sutis são evidentes no tronco encefálico em uma pessoa que está disposta a aceitar o ruído em comparação com aqueles que não estão dispostos a aceitar o ruído.

© 2016 Associação Brasileira de Otorrinolaringologia e Cirurgia Cérvico-Facial. Publicado por Elsevier Editora Ltda. Este é um artigo Open Access sob uma licença CC BY (<http://creativecommons.org/licenses/by/4.0/>).

Introdução

O aparelho auditivo é uma das medidas comuns de reabilitação para pessoas com deficiência auditiva permanente. Em alguns casos de perdas auditivas, ele pode ser usado de maneira transitória. No entanto, os usuários muitas vezes se queixam do ruído de fundo, o que resulta na rejeição.¹ Kochkin² relatou que 31% dos usuários rejeitam o aparelho auditivo por causa do ruído de fundo. Várias medidas de desfecho, que consideram o ruído de fundo como um dos fatores que têm efeito sobre a satisfação com o aparelho auditivo, estão disponíveis. Infelizmente, essas medidas de desfecho devem ser administradas depois de um período de experiência com o aparelho. Além disso, medidas como a fala no teste de ruído, fala rápida no teste de ruído, teste de ruído concorrente e audição em teste de ruído têm sido usadas para prever o seu benefício.³ Embora esses testes sejam sensíveis para medir o desempenho da fala no ruído e administrados no momento do ajuste do aparelho auditivo, eles não conseguem prever o benefício e/ou a satisfação dos aparelhos auditivos na prática clínica.⁴ Esse problema é parcialmente abordado com a mensuração do nível aceitável de ruído (NAR) introduzida por Nabelek et al.,⁵ na qual o cliente classifica o incômodo decorrente do ruído de fundo na presença de fala.

Nabelek et al.⁶ demonstraram que o valor de NAR consegue prever bons e maus usuários de aparelho auditivo com precisão de 85%. O NAR não é afetado pelo tipo de ruído de fundo,⁵ preferência de sons de fundo,⁷ idioma principal

do ouvinte,⁸ níveis de apresentação da fala,⁷ idade, sensibilidade auditiva e teor da língua⁹ ou pelo sinal de fala e sexo do falante.¹⁰ Harkrider¹¹ estudou a correlação fisiológica do NAR envolvido nos centros auditivos superiores e usou a medição eletrofisiológica. Em indivíduos com NAR baixos (ou seja, maior aceitação de ruído de fundo), as amplitudes de onda V de resposta auditiva do tronco encefálico (ABR), os componentes da resposta de média latência (RML) e a resposta de latência tardia (RLT) foram observados como significativamente prolongados, quando comparados com os indivíduos que obtiveram NAR elevados (aceitação de ruído de fundo mais baixa). Isso se deve ao mecanismo eferente mais forte, de tal modo que as entradas sensoriais são suprimidas e/ou o mecanismo aferente central é menos ativo.¹² Assim, demonstrou-se que o NAR é uma medida fisiologicamente sensível. No entanto, é interessante conhecer a maneira pela qual a fala amplificada é representada fisiologicamente em bons e maus usuários de aparelhos auditivos.

Apesar do avanço na tecnologia de aparelhos auditivos, alguns indivíduos os aceitam e outros rejeitam, apesar de a perda auditiva ser semelhante em termos de grau, tipo e configuração. A variabilidade na satisfação de um dispositivo de reabilitação pode ser decorrente dos parâmetros de processamento do aparelho auditivo e/ou da interação entre a potência do aparelho e seus parâmetros acústicos transmitidos através de diferentes partes da via auditiva.¹³ No presente estudo, a produção do aparelho auditivo é investigada no meato acústico externo e no tronco encefálico auditivo.

Há décadas os pesquisadores usam o *Elfitema Probe Tube Microphone* (PTM) para medir o efeito do processamento do aparelho auditivo sobre a acústica da fala. A medição do PTM reflete o efeito acústico de fatores, tais como pavilhão auricular, meato acústico, cabeça e tronco.¹⁴ Primariamente, o PTM é usado para aprimorar/verificar o ganho do aparelho auditivo para combinar com o ganho-alvo em frequências diferentes, conforme prescrito pela fórmula de ajuste.¹⁵ Está bem estabelecido que a saída do aparelho no meato acústico irá alterar amplitudes de formantes que levam à má interpretação. Um experimento foi conduzido por Stelmachowicz et al.,¹⁶ que registraram a saída do aparelho auditivo no meato com uso de aparelhos auditivos lineares e não lineares em três ouvintes com perda auditiva neurosensorial leve a moderada. Eles fizeram análise espectral sobre esses estímulos registrados. Os resultados revelaram uma vazão precipitada na resposta de alta frequência, que limitou as informações de exemplos de consoantes. Em uma linha semelhante, Souza e Tremblay¹⁷ fizeram um estudo para correlacionar erros de consoantes com a análise acústica da fala amplificada em indivíduos com perda auditiva neurosensorial leve a moderada. Eles observaram que o estímulo/da/ foi consistentemente mal percebido como /ga/. Esse fato foi atribuído à amplitude do espectro explosivo auxiliar de/da/, que se verificou como semelhante à amplitude do espectro explosivo não processado de/ga/. Kewley-Port¹⁸ relatou que a identificação de consoantes oclusivas na posição inicial requer o espectro de explosão como o exemplo primário para reconhecimento da fala. Assim, após a amplificação, as consoantes oclusivas são mais propensas a apresentar erros de local.¹⁷ No entanto, a combinação amplificada consoante-vogal da fricativa ou africativa tende a apresentar erros de modo, à medida que a má percepção consistente de sons da fala/ʒi/por/dʒi/¹⁹ for percebida. Quando a produção acústica do aparelho auditivo foi analisada, revelou-se que o espectro de amplitude da fricativa/ʒi/era semelhante à amplitude do espectro não processado da africativa/dʒi/. Assim, fazer análise espectral da produção do aparelho auditivo registrada no meato acústico lança luz sobre os parâmetros de processamento do aparelho auditivo. Há casos em que pistas acústicas são distorcidas, mas um ouvinte ainda reconhece corretamente. Isso pode ocorrer devido à redundância ou em decorrência de pistas contextuais da fala. Em alguns outros casos pistas acústicas são preservadas, mas um ouvinte não consegue reconhecer sons da fala. Possivelmente, isso pode ocorrer devido à sensibilidade insuficiente na cóclea e/ou mudanças concomitantes nos diferentes níveis da via auditiva. Assim, um potencial evocado registrado para estímulos da fala deve ser usado para validar a percepção registrada em diferentes níveis da via auditiva. No presente estudo, a resposta evocada de tronco encefálico auditivo de bons e maus usuários de aparelhos auditivos é investigada.

As frequências seguidas de resposta (FSR) têm sido extensivamente estudadas para se compreender o processamento fisiológico da fala no nível do tronco encefálico auditivo. A FSR é uma resposta de captura de fase a aspectos periódicos de estímulos, inclusive a fala até 1.000 Hz^{20,21} da população neural do cóliculo inferior do tronco encefálico rostral.²² A FSR foi registrada de maneira confiável para sons consoante-vogal/da/.²³⁻²⁶ Além disso, FSR para

estímulo/da/foi investigada em condições monaurais²⁷ e binaurais;²⁸ na presença de ruído de fundo;²⁷ e estimulação da orelha direita ou esquerda.²⁹ A FSR foi registrada com sucesso com alto-falante como transdutor para distribuir os estímulos/da/e/si/.³⁰ A partir daí, é evidente que FSR é uma resposta contingente ao estímulo que é mais consistente para frequências médias e baixas. Embora a resposta de frequência do aparelho auditivo seja de até 6.500 Hz, a FSR é uma ferramenta sensível a qualquer alteração no processamento do tronco encefálico auditivo, de tal maneira que responde à pergunta sobre como os sons da fala amplificados são codificados em frequências médias e baixas de bons e maus usuários de aparelhos auditivos.

A partir da literatura existente, pode-se inferir que as análises espectrais da produção do aparelho auditivo obtidas com o PTM dão informações sobre o processamento do aparelho auditivo no nível do meato acústico. Além disso, o estímulo do meato acústico é transmitido para o nível do tronco encefálico auditivo e é medido com FSR. O FSR vai ajudar a inferir a codificação neural do discurso contínuo. Essas medidas fornecem *insights* sobre a maneira pela qual a fala é neuralmente codificada no nível do tronco encefálico, em indivíduos com deficiência auditiva neurosensorial, que são classificados como bons e maus usuários de aparelhos auditivos, têm tipo, grau e configuração da perda auditiva comparáveis. Portanto, o presente estudo pretendeu investigar a produção do aparelho no meato acústico para determinar a extensão da alteração causada por ele nos parâmetros espectrais. Além disso, a maneira pela qual a fala amplificada é representada no nível do tronco encefálico em bons e maus usuários de aparelhos auditivos também tem sido investigada. Os objetivos formulados para o estudo foram comparar: 1) mudanças espectrais entre BUAA e MUAA em condições com e sem aparelho no meato acústico com PTM; e 2) codificação neural dos sons da fala no nível do tronco encefálico no BUAA e MUAA.

Método

Participantes

Foram selecionados 60 pacientes que apresentavam perda auditiva neurosensorial bilateral moderada com configuração plana. A configuração plana foi operacionalmente definida como a diferença entre os limiares condutivos mais e menos elevados inferiores a 20 dB no intervalo de 0,25 a 8 kHz.³¹ A faixa dos participantes foi de 15 a 65 anos. Eles tinham escores de identificação da fala (EIF) maiores ou iguais a 75% a 40 dB SL (re: limiar de recepção da fala, LRF). A orelha de teste tinha estado normal da orelha média, como indicado por timpanograma tipo 'A' com pressão de pico da orelha média que variou de 50 daPa a -100 daPa, e a entrada variou de 0,5 mL a 1,75 mL. A resposta do tronco encefálico auditivo (ABR) foi gravada em duas taxas de repetição de 11,1 segundos e 90,1 segundos a 90 dB nHL para garantir que não havia doença retrococlear. Observou-se que a diferença de latência de pico V de ABR é inferior a 0,8 ms entre as duas taxas de repetição. Todos os participantes eram usuários iniciais de aparelhos auditivos e não havia história autorrelatada de outros problemas otológicos e neurológicos. Os participantes foram classificados em bons

Tabela 1 Dados demográficos dos participantes do estudo

NE	Idade (anos)	Sexo	Orelha	ATP dB HL	EIF (máx, 25)	Diferença de latência de pico V em duas velocidades de repetição (ms)	Agrupamento baseado em NAR	Escore de NAR (em dB)
1	15,00	F	D	58,30	20,00	0,38	BUAA	3,00
2	18,00	M	E	45,00	21,00	0,44	BUAA	6,00
3	18,00	F	D	56,60	20,00	0,35	BUAA	6,00
4	19,00	M	D	48,30	21,00	0,46	BUAA	6,00
5	19,00	M	E	53,30	25,00	0,27	BUAA	3,00
6	25,00	F	D	55,00	21,00	0,46	MUAA	14,00
7	27,00	M	E	50,00	19,00	0,65	MUAA	15,00
8	27,00	F	D	51,60	19,00	0,55	MUAA	15,00
9	27,00	M	D	51,60	21,00	0,25	MUAA	16,00
10	32,00	M	D	53,30	20,00	0,41	BUAA	6,00
11	32,00	M	D	55,30	22,00	0,33	MUAA	17,00
12	35,00	M	D	51,60	23,00	0,56	BUAA	6,00
13	26,00	F	D	55,00	22,00	0,30	BUAA	6,00
14	37,00	M	E	51,60	21,00	0,34	MUAA	14,00
15	40,00	F	E	50,00	20,00	0,30	BUAA	6,00
16	41,00	M	E	50,00	20,00	0,27	MUAA	17,00
17	42,00	M	E	50,00	25,00	0,37	MUAA	14,00
18	43,00	M	E	55,00	20,00	0,25	BUAA	6,00
19	47,00	F	D	55,60	23,00	0,38	BUAA	4,00
20	47,00	M	D	46,60	19,00	0,37	MUAA	14,00
21	49,00	M	D	53,30	20,00	0,38	BUAA	6,00
22	50,00	F	D	53,30	23,00	0,30	MUAA	14,00
23	50,00	M	E	45,00	25,00	0,55	BUAA	5,00
24	53,00	M	E	45,00	21,00	0,22	MUAA	13,00
25	54,00	F	D	51,60	24,00	0,39	MUAA	14,00
26	54,00	M	D	48,30	24,00	0,28	BUAA	5,00
27	54,00	M	E	41,00	23,00	0,30	BUAA	2,00
28	54,00	M	E	41,60	23,00	0,57	BUAA	4,00
29	55,00	F	E	55,00	21,00	0,24	MUAA	18,00
30	55,00	F	D	53,30	23,00	0,50	BUAA	3,00
31	55,00	F	E	50,00	23,00	0,58	BUAA	2,00
32	51,00	M	E	45,00	24,00	0,50	BUAA	6,00
33	56,00	F	E	46,60	23,00	0,24	BUAA	4,00
34	58,00	M	D	45,00	24,00	0,34	BUAA	6,00
35	58,00	M	E	48,30	24,00	0,32	BUAA	6,00
36	58,00	M	D	45,00	24,00	0,24	BUAA	4,00
37	58,00	M	E	45,00	24,00	0,34	BUAA	2,00
38	60,00	F	D	41,60	21,00	0,32	BUAA	6,00
39	60,00	F	E	53,30	21,00	0,22	BUAA	2,00
40	60,00	F	D	45,00	24,00	0,54	MUAA	15,00
41	60,00	M	D	55,00	20,00	0,43	MUAA	14,00
42	60,00	M	E	43,30	25,00	0,55	BUAA	3,00
43	60,00	M	D	46,30	24,00	0,37	MUAA	15,00
44	60,00	M	E	46,30	24,00	0,23	MUAA	14,00
45	60,00	M	D	58,30	20,00	0,34	MUAA	16,00
46	61,00	F	E	55,00	23,00	0,44	BUAA	5,00
47	61,00	F	D	58,30	24,00	0,36	BUAA	6,00
48	61,00	M	D	41,60	24,00	0,70	MUAA	16,00
49	61,00	M	D	55,00	21,00	0,42	MUAA	14,00
50	62,00	M	D	55,00	19,00	0,26	BUAA	2,00
51	62,00	M	E	51,60	25,00	0,28	MUAA	13,00
52	62,00	F	D	50,00	24,00	0,34	MUAA	14,00
53	64,00	M	E	45,00	20,00	0,23	MUAA	16,00
54	64,00	F	E	55,30	22,00	0,28	BUAA	5,00
55	65,00	M	E	51,60	21,50	0,28	MUAA	14,00
56	65,00	M	D	55,00	24,00	0,32	BUAA	5,00
57	65,00	M	D	46,60	21,50	0,30	BUAA	5,00
58	65,00	M	D	43,30	20,00	0,35	BUAA	6,00
59	65,00	M	D	50,00	18,00	0,68	BUAA	6,00
60	65,00	F	D	56,60	24,00	0,45	MUAA	14,00

ATP, audiometria tonal pura; BUAA, bons usuários de aparelho auditivo; D, orelha direita; E, orelha esquerda; EIF, escores de identificação da fala; F, feminino; M, masculino; MUAA, maus usuários de aparelho auditivo; NAR, nível aceitável de ruído; PA, perda auditiva.

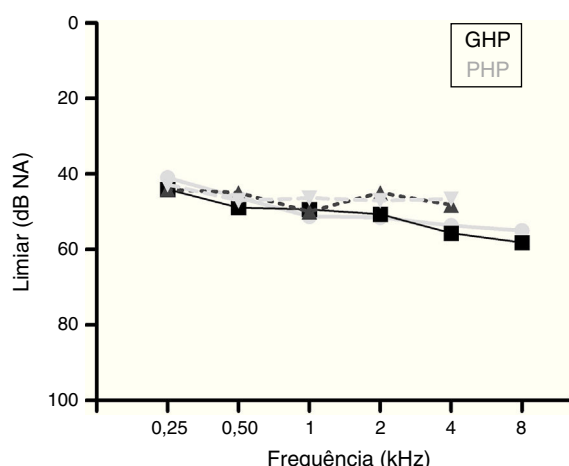


Figura 1 Audiogramas de bons e maus usuários de aparelhos auditivos.

ou maus usuários de aparelho auditivo (BUAA ou MUAA) por meio do uso do nível aceitável de ruído (NAR).⁶ Os participantes que obtiveram um escore de NAR de ≤ 7 foram considerados como bons usuários de aparelhos auditivos e com um escore ≥ 13 foram considerados como maus usuários de aparelhos auditivos.⁶ Os dados demográficos de cada participante no grupo clínico são apresentados na [tabela 1](#). Os limiares auditivos em cada frequência audiométrica da orelha de teste dos bons e maus usuários de aparelhos auditivos estão representados na [figura 1](#). O estudo foi aprovado pelo Comitê de Ética para pesquisa em seres humanos do *All India Institute of Speech and Hearing*. O consentimento informado foi obtido de cada participante.

Estímulos

Tokens consoante-vogal (CV) produzidos naturalmente foram usados como estímulos de teste. Um homem adulto com voz normal foi usado para gravar os tokens de CV. As durações dos estímulos/da/e/si/ foram de 94 ms e 301 ms, respectivamente. Para/da/, o tempo de início da sonorização foi de 18 ms, a duração da explosão de 5 ms, a duração da transição de 37 ms e a duração da vogal de 34 ms. Para/si/, a duração fricativa foi de 159,3 ms, duração da transição de 47,1 ms e duração da vogal de 94,6 ms. Ambos os estímulos foram convertidos a partir de formato '.wav' para '.avg' com wavtoavg de m-file da Brainstem toolbox. O formato '.avg' de ambos os estímulos foi filtrados com

filtro passa-faixa de 30-3000 Hz com Neuroscan (Scan 2-versão 4.4) para saber a relação funcional entre a estrutura acústica da fala e a resposta do tronco encefálico à fala. As formas de onda 'stimulus.avg' e espectrogramas dos dois tokens de CV estão representados na [figura 2](#). A [tabela 2](#) resume a frequência fundamental (F_0) e as duas primeiras frequências formantes (F_1 e F_2) dos estímulos do componente de vogal/da/e/si/. O início do estado estacionário F_0 , F_1 e F_2 dentro do período de transição (37 ms) do estímulo de/da/ e componentes de frequência dentro da duração da transição (42 ms) da porção/i/do estímulo de/si/ foram medidos com o *software* Praat (versão 5.1.29).

Além disso, a passagem Kannada desenvolvida por Sairam e Manjula³² foi lida em voz alta com esforço vocal normal por uma falante do sexo feminino e gravada com *software* Adobe Audition (versão 3). Essa passagem registrada foi usada para determinar o nível aceitável de ruído (NAR). Um teste foi feito para verificar a qualidade da passagem Kannada gravada, no qual dez ouvintes com audição normal avaliaram a passagem para naturalidade.

Aparelho auditivo

O aparelho auditivo digital retroauricular (BTE) foi usado para gravar a saída no meato acústico e na resposta auditiva do tronco encefálico de cada participante. De acordo com as especificações técnicas, a gama de frequência do aparelho auditivo teste variou de 0,210 a 6,5 kHz. O ganho completo de pico foi de 58 dB e o ganho completo médio de alta frequência foi de 49 dB. O funcionamento do aparelho auditivo foi assegurado no início da coleta de dados e repetido periodicamente durante a coleta de dados.

Procedimento

Cada participante foi classificado como bom e mau usuário de aparelho auditivo por meio do teste NAR comportamental. O aparelho auditivo teste com molde auricular personalizado foi ajustado em cada participante e o seu ganho foi aprimorado. Para melhorar o resultado do aparelho auditivo, seis sílabas de Ling foram apresentadas em um nível calibrado de 65 dB NPS por meio do audiômetro em um campo sonoro. O ganho e a resposta de frequência do aparelho auditivo foram manipulados para audibilidade de cada uma das seis sílabas de Ling, por meio da opção de ajuste fino. Para saber em que medida a característica espectral é preservada pelo aparelho auditivo, a produção do

Tabela 2 Frequência fundamental e as duas frequências formantes (em Hz) na duração da transição da versão original e filtrada de estímulos/da/e/si/

Estímulo		F0 (Hz)		F1 (Hz)v		F2 (Hz)	
		Início	Estado estacionário	Início	Estado estacionário	Início	Estado estacionário
Versão original	/da/	135,7	131,2	519,8	556,3	1.822,4	1.677,7
	/si/	145,7	137,5	345,4	308,8	2.268,5	2.451,5
Versão filtrada	/da/	135,7	131,2	519,8	556,3		
	/si/	145,7	137,5	345,4	308,8		

F_0 , frequência fundamental; F_1 e F_2 , primeira e segunda frequências formantes; F_2 para versão filtrada não é aplicável, pois a frequência de corte máxima do filtro foi 2 kHz.

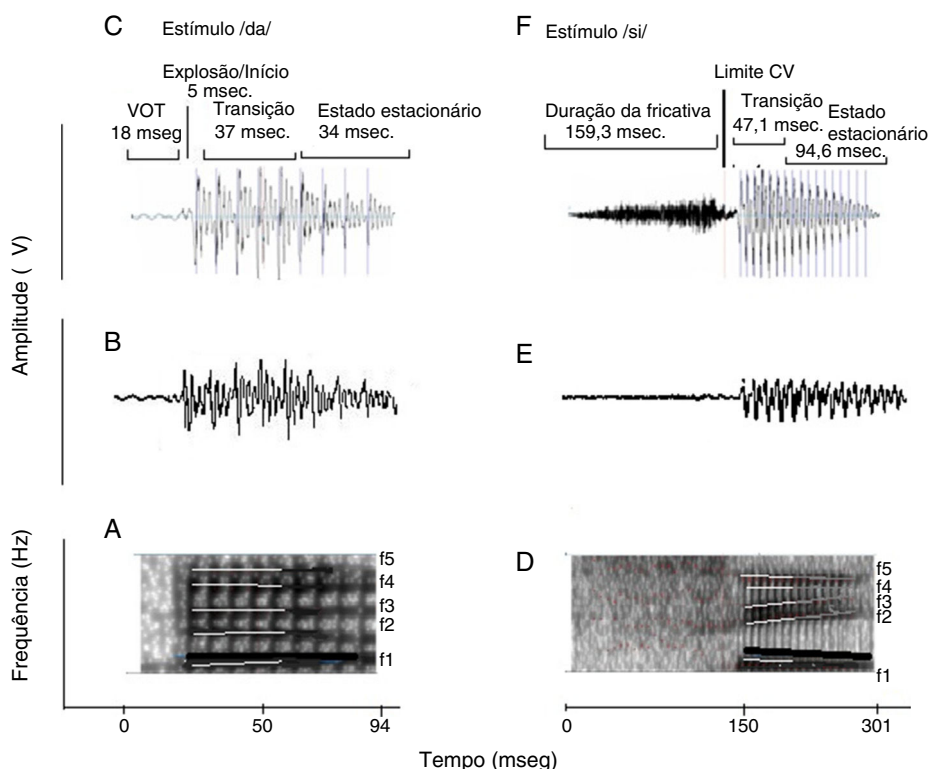


Figura 2 A) e D) são os espectrogramas de estímulos/da/e/si/. A linha sólida preta em ambos os estímulos indica a F_0 , que tem um padrão de queda. As frequências formantes são representadas por linhas brancas. A F_1 e F_2 do/da/ tem padrão plano. A F_1 de estímulo/si/ tem padrão de queda e F_2 tem padrão de elevação; C) e F) são formas de onda dos estímulos/da/e/si/. Para som/da/, o tempo de início de voz foi de 18 ms, a duração da explosão de 5 ms, a duração da transição de 37 ms e a duração da vogal de 34 ms. Para o som/si/, a duração da fricativa foi de 159,3 ms, duração da transição de 47,1 ms e duração da vogal de 94,6 ms.

aparelho para cada estímulo foi gravada no meato acústico com a medida do microfone sonda. Além disso, a FSR no nível do tronco encefálico foi gravada para cada estímulo, em ambas as condições, com e sem aparelho.

Nível aceitável de ruído

O nível aceitável de ruído (NAR) avalia a reação do ouvinte ao ruído de fundo enquanto ouve uma fala. Para a medição do NAR, o método fornecido por Nabelek et al.⁵ foi adotado. Cada participante do estudo foi sentado confortavelmente em uma cadeira, em frente ao alto-falante do audiômetro, que estava a 1 m de distância e 45° de azimuth. Para calcular o NAR, o nível de mais conforto (NMC) e o nível de ruído de fundo (NRF) foram medidos.

A passagem Kannada registrada foi acompanhada desde a entrada auxiliar até o alto-falante do audiômetro. O nível de apresentação foi definido no nível de SRT. Aos poucos, ele foi ajustado em etapas de 5 dB para estabelecer o nível mais confortável (NMC) e, em seguida, em etapas de menor tamanho de +1 e -2 dB, até que o NMC foi estabelecido de maneira confiável. Após o NMC ser estabelecido, o ruído na fala foi introduzido a 30 dB NA. O grau desse ruído foi aumentado em etapas de 5 dB inicialmente e depois em etapas de 2 dB, até um ponto em que o participante estivesse disposto a aceitar o ruído sem sentir cansaço ou fadiga enquanto ouvia e acompanhava as palavras da história. O grau máximo que ele poderia aceitar ou tolerar com o ruído na fala sem se

cansar foi considerado como o nível de ruído de fundo (NRF). O nível de ruído na fala foi ajustado até o participante ser capaz de suportar o ruído enquanto acompanhava a história. O grau resultante foi o NRF. O NAR quantifica o nível aceitável de ruído de fundo e é calculado como a diferença entre NMC (dB NA) e NRF (dB NA).⁴ Com base nos NAR, cada participante foi classificado como bom (NAR de ≤ 7 dB) ou mau (NAR de ≥ 13 dB) usuário do aparelho auditivo.⁴ O procedimento de NAR foi repetido duas vezes e a média dos dois valores foi considerada como o NAR para cada participante.

Aprimoramento do ganho com aparelho auditivo

Cada participante foi equipado com o aparelho auditivo teste BTE digital personalizado com um molde de proteção macio. O aparelho foi programado com uma fórmula prescritiva de NAL -NL1. A medição real da orelha foi feita para combinar objetivamente o ganho do aparelho auditivo com o estabelecido como alvo. Além disso, seis sons da fala de Ling foram apresentados a 65 dB NPS para aprimorar o ganho do aparelho auditivo. Por meio da opção de ajuste fino, o ganho e a frequência do aparelho auditivo foram aprimorados para a audibilidade dos seis sons de Ling.

Fala processada do aparelho no meato acústico

O nível de cada sinal (armazenado no computador pessoal) foi variado no audiômetro de modo que a intensidade

medida foi de 65 dB NPS no medidor de nível sonoro. O medidor de nível sonoro (MNS) da marca Larson Davis 824 foi posicionado na orelha teste do participante. O MNS foi fixado em função de ponderação rápida e procurou-se confirmar se os estímulos de/da/e/si/ foram apresentados a 65 dB NPS, com base no nível de amplitude de pico lido no MNS. Após a confirmação da calibração dos estímulos, o espectro de saída no meato acústico foi registrado com a medição com microfone sonda, tanto com como sem o aparelho auditivo. O microfone sonda no meato acústico capta as energias espectrais em bandas de aproximadamente meia oitava de 0,25 kHz a 8 kHz para cada estímulo de fala. Os níveis, como função da frequência de 0,25 kHz a 8 kHz, em oitavas, foram anotados para cada estímulo, em condições com e sem aparelho.

Aquisição da frequência seguida de resposta

Cada participante foi confortavelmente sentado em uma cadeira reclinável com braço. Os locais dos eletrodos foram limpos com gel de preparação para a pele. Eletrodos revestidos de prata do tipo disco foram colocados com gel de condução nos locais de teste. A FSR foi gravada com montagem vertical. O eletrodo não inversor (+) foi posicionado no vértice (Cz), o eletrodo terra na parte superior da testa (Fpz) e o eletrodo de inversão (-) no nariz. Foi assegurado que a impedância do eletrodo era inferior a 5 K ohms para cada um dos eletrodos e que a impedância entre eletrodos era inferior a 2 K ohms.

Antes da gravação, a calibração dos estímulos foi confirmada com o MNS do Sistema Larson Davis 824. O MNS foi posicionado no ponto de referência, o local onde a orelha de teste do participante seria posicionada no momento da avaliação. A MNS foi estabelecida em função de ponderação rápida. Certificou-se de que ambos os estímulos de/da/e/si/ foram apresentados a 65 dB NPS, com base no nível de amplitude de pico lido no MNS.

O alto-falante do equipamento de potencial evocado auditivo foi colocado a 45° azimute a partir da orelha teste do participante, na posição calibrada de 12 polegadas de distância. A altura do alto-falante foi ajustada para o nível da orelha do participante. O participante foi instruído a ignorar o estímulo e assistir a um filme, que foi silenciado e rodado em um *notebook* operado por bateria. Também foi solicitado ao participante que reduzisse o movimento dos olhos e da cabeça.

Para registrar a FSR com e sem aparelho, o estímulo de/da/foi apresentado com o alto-falante no nível de apresentação de 65 dB NPS para a orelha teste. O sistema de potencial evocado com base em um PC, Neuroscan 4.4 (Stim 2-versão 4.4), controlou o tempo de apresentação do estímulo e distribuiu um gatilho externo ao sistema de gravação do potencial evocado, Neuroscan (Scan 2-versão 4.4). Para possibilitar um período refratário suficiente dentro da varredura do estímulo, enquanto era reduzido o tempo total de gravação, um intervalo interestímulos (ISI) de 93 ms desde a compensação até o início do próximo estímulo foi incluído para gravar FSR para estímulo de/da/. Um procedimento semelhante foi repetido para gravar a FSR com e sem aparelho para o estímulo de/si/. No entanto, para a gravação desse FSR, usou-se um ISI de 113 ms. A ordem de estímulos ao

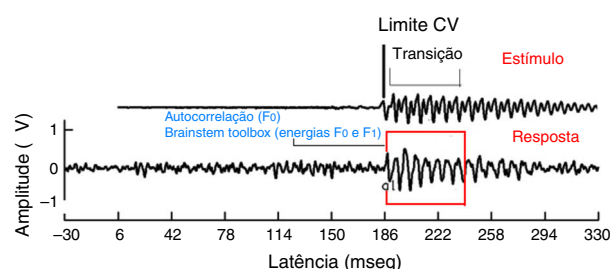


Figura 3 Resposta de transição da FSR obtida de estímulo de/si/ em condição com aparelho.

testar cada participante foi contrabalançada. A FSR foi gravada a partir de 1.500 varreduras cada em polaridades de condensação e rarefação, distribuídas em um trem homogêneo com o *software* de apresentação do estímulo Neuroscan 4.4 (Stim 2-versão 4.4).

A gravação da FSR foi iniciada após a obtenção de um eletroencefalograma (EEG) estável. O EEG contínuo foi convertido de analógico para digital com a taxa de 20.000 Hz. O EEG contínuo foi filtrado on-line com filtro passa-faixa de 30 a 3000 Hz com roll-off de 12 dB/oitava. Posteriormente, foi armazenado em disco para análise *off-line*.

Análise dos dados

A produção do aparelho de audição no meato acústico para cada estímulo nas condições com e sem aparelho foi analisada para espectros. Além disso, a FSR registrada foi analisada para F_0 , energia F_0 e energia F_1 obtidas para cada estímulo. Os dados de EEG contínuos foram processados (*epoch*) sobre uma janela de 160 ms para estímulo de/da/ (que incluiu um período pré-estímulo de 30 ms e um tempo pós-estímulo de 130 ms). A resposta para o estímulo de/si/ foi processada sobre uma janela de 360 ms (que incluiu um período pré-estímulo de 30 ms e um período pós-estímulo de 330 ms). As formas de onda processadas (*epoch*) foram corrigidas para momento basal. As respostas foram ponderadas e filtradas *off-line* a partir de 0,030 kHz (filtro passa-alta, 12 dB/oitava) a 3 kHz (filtro passa-baixa, 12 dB/oitava). Todos os artefatos acima de ± 35 microvolts foram rejeitados, enquanto era feita a média da resposta para cada resposta média, em polaridade de rarefação e condensação. Um mínimo de 1.450 épocas livres de artefato foi assegurado. As formas de onda médias de polaridade de rarefação e condensação foram adicionadas. Além disso, as formas de onda adicionadas foram criadas ao se fazer a média de dois ensaios gravados para cada estímulo, em condições sem aparelho e com aparelho.

Para todos os participantes, as respostas sem aparelho estavam ausentes para ambos os estímulos. Da FSR registrada para o estímulo de/da/com aparelho, a latência do pico 'V' foi identificada por inspeção visual. O padrão MATLAB-código de autocorrelação foi usado, no qual uma faixa de latência foi especificada para obter F_0 na FSR, isto é, da latência do pico 'V' observado até a duração de transição. Enquanto isso, a latência de 'a1' correspondente ao limite CV³⁰ na FSR foi identificada para o estímulo de/si/ e é apresentada na [figura 3](#). A latência de 'a1' até a duração da transição foi especificada no código de autocorrelação

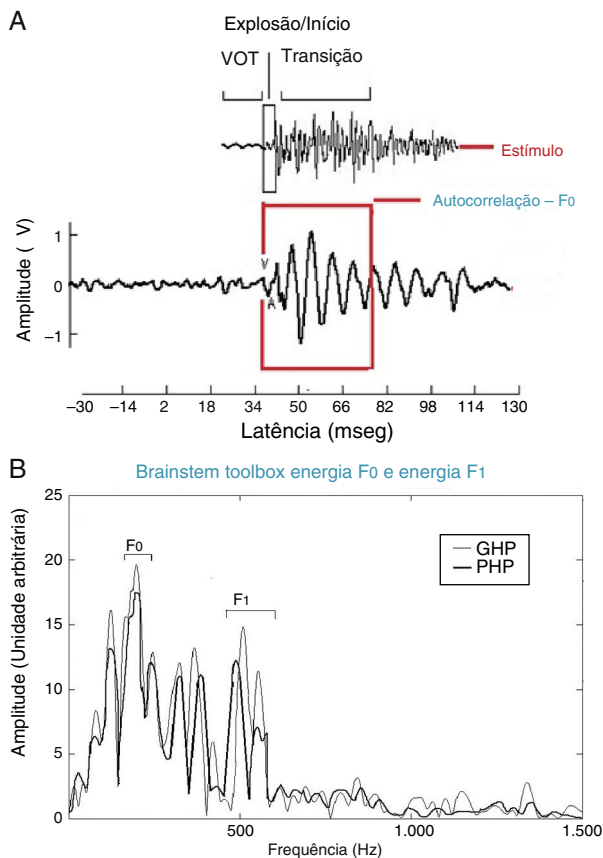


Figura 4 A) Resposta que corresponde ao estímulo na porção de transição do estímulo/da/; B) Mostra grande espectro médio de subgrupo BUAA e MUAA. B) Frequência fundamental e frequência do primeiro formante em FSR para estímulo/da/.

MATLAB para obter F_0 na FSR. Além disso, foram determinadas energia F_0 e energia F_1 , com Brainstem toolbox que usa a técnica FFT (fig. 4), a partir da resposta transitória (pico 'V' para estímulo/da/; e 'a1' para estímulo/si/) até duração da transição especificada (37 ms para estímulo/da/; e 47,1 ms para estímulo/si/).³³

Resultados

Os dados espectrais obtidos no meato acústico com uso da medição com sonda e FSR no nível do tronco encefálico foram analisados em bons e maus usuários do aparelho auditivo. O *software* Statistical Package for the Social Sciences (SPSS para Windows, versão 17) foi usado para fazer as análises estatísticas. Os resultados obtidos foram discutidos com relação a cada objetivo.

Produção da fala do aparelho auditivo no meato acústico

A energia espectral em frequências de 0,25 a 8 kHz (em oitavas) com e sem aparelhos, para ambos os estímulos, foi analisada. Foi feita para determinar a representação de energia por meio de frequências no meato acústico, em bons e maus usuários de aparelho auditivo. Os dados de energia espectral satisfizeram a suposição de distribuição

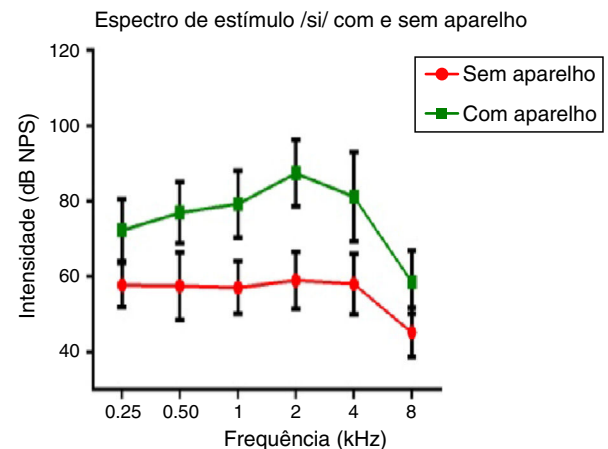


Figura 5 Desvio médio e padrão de intensidade de estímulo/si/ como função da frequência em condições com e sem aparelho.

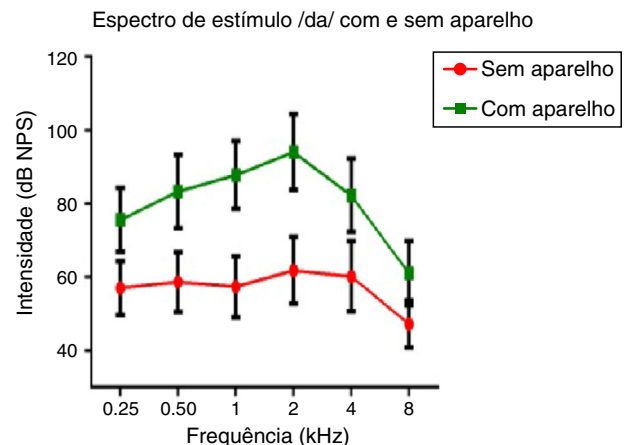


Figura 6 Desvio médio e padrão de intensidade de estímulo/da/ como função da frequência em condições com e sem aparelho.

normal no teste de normalidade de Kolmogorov-Smirnov ($p > 0,05$) e homogeneidade no teste de Levene ($F < 2$). A energia espectral (0,5 kHz a 8 kHz em oitavas) para ambos os estímulos obtidos de ambos os grupos, em condições com e sem aparelhos, foi submetida a Manova. O resultado revelou que não havia diferença significativa entre os grupos na energia espectral em cada frequência de oitava, tanto em condições com aparelho como sem aparelho, para estímulos/da/e/si/. Assim, os dados de energia espectral foram combinados entre os grupos. A análise descritiva foi feita separadamente nas condições sem aparelho e com aparelho. Para estímulo/da/ (fig. 5), no extremo de baixa frequência (0,25 kHz) e em frequências extremamente altas (4 kHz e 8 kHz), a energia em ambas as condições, com e sem aparelho, é relativamente mínima em relação a outras (0,5 kHz, 1 kHz e 2 kHz). Para estímulo/si/ (fig. 6), a baixas frequências extremas (0,25 kHz) e em alta frequência extrema (8 kHz), a energia em ambas as condições, com e sem aparelho, é relativamente mínima, em comparação com outras frequências (1 kHz, 2 kHz e 4 kHz).

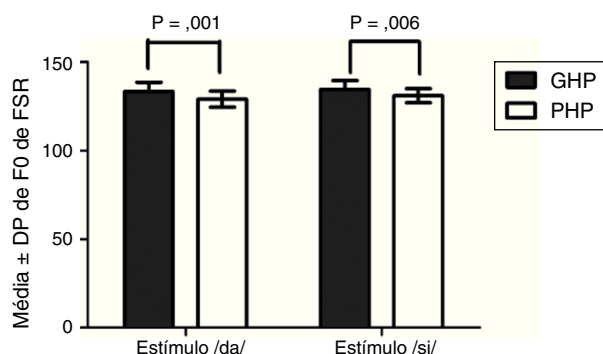


Figura 7 Desvio médio, padrão e valor p de teste t de amostras independentes em F_0 de FSR para cada estímulo, nos grupos BUAA e MUAA.

Comparação de FSR em termos de F_0 , energia F_0 e energia F_1 em bons e maus usuários de aparelhos auditivos

A F_0 , a energia F_0 e a energia F_1 de FSR entre os grupos para cada estímulo satisfizeram a suposição de distribuição normal no teste de normalidade de Kolmogorov-Smirnov ($p > 0,05$) e o teste de homogeneidade de Levene ($F < 2$) também foi feito. Assim, foi feito um teste t de amostras independentes em cada um dos dados de FSR entre os grupos de BUAA ($n = 34$) e MUAA ($n = 24$). A partir do valor médio de F_0 de FSR (fig. 7), pode-se inferir que F_0 foi mais bem representada em BUAA do que em MUAA, para cada estímulo. Além disso, F_0 da FSR foi comparada entre BUAA e MUAA com o teste de amostras independentes. O resultado mostrou que havia uma F_0 significativamente melhor que codificava em BUAA do que em MUAA para o estímulo/*da*/ ($t = 3,41$, $p = 0,001$) e estímulo/*si*/ ($t = 2,84$, $p = 0,006$).

Além disso, as energias F_0 e F_1 da FSR foram comparadas entre os grupos BUAA e MUAA para cada estímulo. Na figura 8, observou-se que a média e o desvio padrão da energia F_0 e da energia F_1 de FSR para cada estímulo foram maiores em BUAA do que em MUAA. Além disso, para saber se havia alguma diferença significativa entre BUAA e MUAA na média da energia F_0 e a energia F_1 da FSR para cada estímulo, um teste t de amostras independentes foi feito. O resultado revelou uma energia F_0 significativamente maior no BUAA do que em MUAA para estímulo/*da*/ ($t = 6,80$, $p = 0,000$) e estímulo/*si*/ ($t = 6,20$, $p = 0,000$). Além disso, observou-se uma energia F_1 significativamente maior em BUAA do que em MUAA para estímulo/*da*/ ($t = 3,11$, $p = 0,002$) e estímulo/*si*/ ($t = 5,20$, $p = 0,000$).

Discussão

O objetivo do estudo foi investigar a representação da fala amplificada no meato acústico e no tronco encefálico auditivo de bons e maus usuários de aparelhos auditivos.

Efeito do processamento do aparelho auditivo nos parâmetros espectrais de estímulos da fala

Na condição sem aparelho para estímulos/*da*/e/*si*/, a energia medida a 2 kHz foi relativamente maior do que em

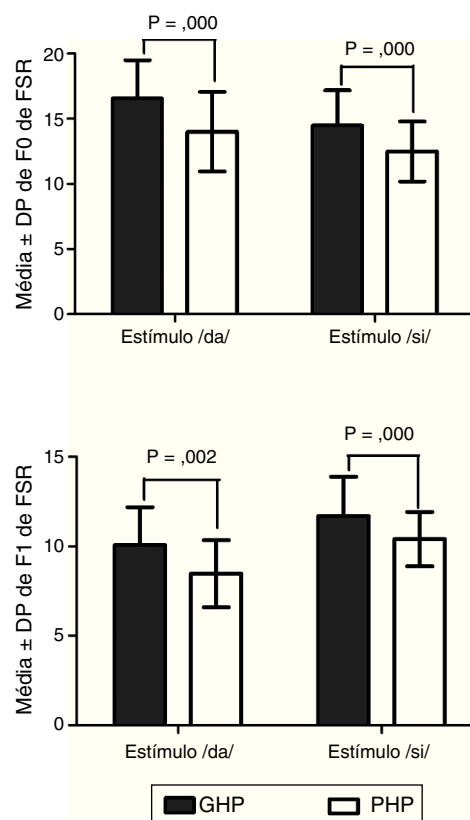


Figura 8 Desvio médio e padrão de energias F_0 e F_1 para cada estímulo em grupos de BUAA e MUAA.

outras frequências (em oitavas). Além disso, houve uma queda na energia depois de 4 kHz, ou seja, cerca de 10 dB por oitava para estímulo/*da*/e 12 dB por oitava para estímulo/*si*/ (figs. 5 e 6). Esse padrão de representação de energia em função da frequência, para ambos os estímulos, poderia ser por causa da resposta de frequência de microfone usada para gravar os estímulos de teste-alvo. Na condição com aparelho para estímulo/*da*/e estímulo/*si*/, a energia medida foi relativamente menor nos dois extremos de frequências de corte, que em baixa frequência fica abaixo de 0,25 kHz e em altas frequências acima de 4 kHz. Assim, com frequências extremas, a energia média em ambas as condições, com e sem aparelho, foi menor. Em baixa frequência, a energia reduzida poderia ser justificada pelo menor ganho naquela região de frequência fornecida pela fórmula prescritiva.³⁴ Além disso, a baixa energia observada na região de baixa frequência de estímulo/*da*/e/*si*/ também ocorreria devido à resposta de frequência do aparelho auditivo. O corte da baixa frequência da resposta de frequência do aparelho auditivo de teste foi de 0,21 kHz. Em altas frequências, ou seja, acima de 4 kHz, a energia reduziu aproximadamente à taxa de 10 dB por oitava para estímulos/*da*/e 14 dB por oitava para estímulo/*si*/ . Essa poderia ser a resposta de frequência dos estímulos/*da*/e/*si*/ que tinham energia até 4 kHz, como observado na condição sem aparelho. No entanto, outra justificativa pode se dever ao fato de que embora a resposta de frequência do aparelho auditivo tivesse 0,216 a 6,5 kHz, a energia após 4 kHz gradualmente reduziu por oitava. Assim, observou-se energia notável na gama de frequências de 0,5 kHz a 4 kHz. Pode-se inferir

que há uma amplificação relativamente mais alta na região de frequência média do aparelho auditivo do que outras duas frequências extremas de corte (baixa e alta). Informalmente, os participantes foram instruídos a repetir as sílabas apresentadas aleatoriamente por três vezes. Na condição sem aparelho, os participantes foram incapazes de identificar os tokens de CV quando o nível de apresentação foi de 65 dB NPS e não conseguiram chegar à faixa de audibilidade. No entanto, na condição com aparelho, todos os participantes consistentemente identificaram sílabas. Além disso, na análise espectral, observou-se que a amplitude do espectro de explosivas com aparelho de /da/ era semelhante ao espectro da amplitude de explosiva não processada de /da/. Observou-se também que o espectro de amplitude de fricativa /si/ era semelhante à amplitude do espectro de fricativa não processada de /si/. Infere-se que o aparelho auditivo preserva pistas de fala inerentes ao meato acústico.

Comparação de F_0 da FSR, energia F_0 e energia F_1 em bons e maus usuários de aparelhos auditivos

A FSR em ambas as condições com e sem aparelho foi obtida de todos os participantes. Na condição sem aparelho, as respostas do tronco encefálico estavam ausentes, pois os estímulos (/da/e/si/) foram apresentados a 65 dB NPS, que não conseguiram chegar à audibilidade. Na condição com aparelho, a representação na FSR para cada estímulo (/da/e/si/) manteve-se substancial e semelhante à do estímulo bruto filtrado não processado. Isso indicou que o conteúdo espectral preservado do aparelho auditivo é transmitido para o nível de tronco encefálico auditivo. Para estímulo /da/, a F_0 média de FSR foi maior em BUAA (133,46 Hz) do que em MUAA (128,84), de tal modo que se observou que a diferença era significativamente variada. Isso foi verdadeiro para F_0 da FSR para estímulo /si/ entre BUAA (134,42 Hz) e MUAA (130,84 Hz). Além disso, F_0 do estímulo de /da/ com aparelho foi 134,95 Hz e para /si/ foi 144,74 Hz. A diferença em F_0 (em Hz), entre codificação de F_0 no nível do tronco encefálico e F_0 de estímulo de teste com aparelho, foi de 1 Hz em BUAA e 6 Hz em MUAA para estímulo /da/. Da mesma maneira, a diferença observada foi de 10 Hz no BUAA e 14 Hz em MUAA para estímulo /si/.

A diferença média entre BUAA e MUAA na codificação de F_0 foi de 5 Hz para estímulo /da/ e 4 Hz para /si/. Embora essa diferença tenha sido significativa na codificação de F_0 entre BUAA e MUAA para ambos os estímulos, isso pode não trazer uma mudança na identidade do falante. Isso porque, de acordo com Iles³⁵ uma mudança de até ± 25 Hz na F_0 não provoca uma mudança na identidade do falante. O achado do estudo está de acordo com o relatório de pesquisa feito por Horii,³⁶ que relatou que uma diferença superior a 25 Hz na F_0 entre os mesmos dois estímulos não causa diferença de identidade do falante. Além disso, a variabilidade intraindividual de F_0 em um esforço vocal normal variou entre $\pm 9,6$ Hz.³⁷ Assim, pode-se inferir que F_0 média de FSR para estímulos /da/e/si/ foi neuralmente mais bem representada em BUAA do que em MUAA e que ambos os grupos foram capazes de reconhecer a identidade do falante.

Além disso, observou-se que a energia F_0 e a energia F_1 de FSR para cada estímulo foram significativamente maiores em BUAA do que em MUAA. As energias mais elevadas

de F_0 e F_1 em BUAA podem ser causadas por fibras eferentes mais fortes, que inibem outros harmônicos que não correspondem à frequência fundamental e às frequências formantes. Isso está de acordo com os relatos da pesquisa de Ashmore³⁸ e Knight.³⁹ Para ser mais específico, o mecanismo aferente central seria mais forte no grupo de BUAA, de tal modo que os neurônios no colículo inferior disparariam precisamente para os harmônicos correspondentes a F_0 e F_1 . Além disso, o mecanismo eferente pode ser mais forte, de modo que as fibras eferentes inibem outros harmônicos que não correspondem à frequência fundamental e às frequências formantes, ajustam assim a entrada auditiva. Os mecanismos excitatórios e inibitórios de neurônios do gerador neural subjacente do colículo inferior em BUAA dispararam de maneira mais ou menos precisa aos componentes F_0 e F_1 correspondentes do estímulo. A inferência do presente estudo sustenta os achados de Krishnan.⁴⁰ Ele demonstrou que a via auditiva eferente suprime energias adjacentes aos harmônicos correspondentes a F_0 e F_1 da FSR. Juntamente com uma via aferente ativa, o nervo auditivo aferente gera a atividade elétrica e corresponde mais precisamente a F_0 e F_1 do estímulo. Isso envolve a liberação de neurotransmissor, reduzi, assim, o limiar transmembrana e um aumento no disparo neural. Em maus usuários de aparelhos auditivos, embora haja presença de atividade fisiológica semelhante, provavelmente uma falta de precisão na atividade neural devido a aferentes menos sensíveis e via auditiva eferente pode ter falhado em fornecer maior energia em harmônicos correspondentes a F_0 e F_1 de cada estímulo. Assim, pode-se inferir, a partir do presente estudo, que as variações fisiológicas sutis podem estar presentes no colículo inferior da via auditiva nos maus usuários de aparelhos auditivos, com relação àquelas em bons usuários de aparelho auditivo.

Conclusão

Embora o aparelho auditivo tenha preservado pistas inerentes nas sílabas de fala, o desconforto com o ruído altera a codificação neural no nível auditivo do tronco encefálico. O estudo conclui que as pistas acústicas transferidas pelo aparelho auditivo são transmitidas com sucesso no nível de tronco encefálico, mas alterações fisiológicas sutis se encontram presentes no tronco encefálico auditivo nos indivíduos desconfortáveis com ruído, em comparação com aqueles sem desconforto.

Implicação

O estudo apresenta uma evidência do uso de abordagens objetivas para validar a produção do aparelho para audição no meato acústico e no nível do tronco encefálico auditivo. O uso da medição real da orelha para analisar a produção do aparelho no meato acústico ajudará a conhecer a representação de pistas específicas da fala. O estudo da codificação da fala amplificada em indivíduos com deficiência auditiva com o seu nível de desconforto demonstra um papel crítico da resposta de estímulo contingente na avaliação de algoritmos de aparelhos auditivos. Ele resolve alguns dos problemas práticos enfrentados pelos fonoaudiólogos com relação à definição de parâmetros de amplificação

no fornecimento máximo de informação usável. Os achados do presente estudo ajudam o fonoaudiólogo no aconselhamento de um usuário de aparelho auditivo em relação à extensão do benefício obtido com o melhor aparelho auditivo prescrito.

Conflitos de interesse

Os autores declaram não haver conflitos de interesse.

Agradecimentos

Ao diretor e ao HOD-Audiologia, Instituto All India de Fonoaudiologia, por permitirem o estudo. E a todos os participantes pela cooperação.

Referências

- Bentler RA, Niebuhr DP, Getta JP, Anderson CV. Longitudinal study of hearing aid effectiveness. II: Subjective measures. *J Speech Hear Res.* 1993;36:820-31.
- Kochkin S. On the issue of value: hearing aid benefit, price, satisfaction, and brand repurchase rates. *Hear Rev.* 2003;10:12-26.
- Nemes J. Despite benefits of outcomes measures, advocates say they're under used. *Hear J.* 2003;56:19-25.
- Nabelek AK. Acceptance of background noise may be key to successful fittings. *Hear J.* 2005;58:10-5.
- Nabelek AK, Tucker FM, Letowski TR. Tolerant of background noises: relationship with patterns of hearing aid use by elderly persons. *J Speech Hear Res.* 1991;34:679-85.
- Nabelek AK, Freyaldenhoven MC, Tampas JW, Burchfiel SB, Muenchen RA. Acceptable noise level as a predictor of hearing aid use. *J Am Acad Audiol.* 2006;17:626-39.
- Freyaldenhoven MC, Smiley DF. Acceptance of background noise in children with normal hearing. *J Educ Audiol.* 2006;13:27-31.
- von Hapsburg D, Bahng J. Acceptance of background noise levels in bilingual (Korean-English) listeners. *J Am Acad Audiol.* 2006;17:649-58.
- Brannstrom KJ, Lantz J, Nielsen LH, Olsen SO. Acceptable noise level with Danish, Swedish, and non-semantic speech materials. *Int J Audiol.* 2012;51:146-56.
- Plyler PN, Alworth LN, Rossini TP, Mapes KE. Effects of speech signal content and speaker gender on acceptance of noise in listeners with normal hearing. *Int J Audiol.* 2011;50:243-8.
- Harkrider AW, Tampas JW. Differences in responses from the cochleae and central nervous systems of females with low versus high acceptable noise levels. *J Am Acad Audiol.* 2006;17:667-76.
- Tampas JW, Harkrider AW, Hedrick MS. Neurophysiological indices of speech and nonspeech stimulus processing. *J Speech Lang Hear Res.* 2005;48:1147-64.
- Souza PE, Tremblay KL. New perspectives on assessing amplification effects. *Trends Amplif.* 2006;10:119-43.
- Mueller HG. Probe microphone measurement: unplugged. *Hear Rev.* 2005;48:10-36.
- Cunningham DR, Lao-Davila RG, Eisenmenger BA, Lazich RW. Study finds use of Live Speech Mapping reduces follow-up visits and saves money. *J Hear.* 2002;55:43-6.
- Stelmachowicz PG, Kopun J, Mace A, Lewis DE, Nitttrouer S. The perception of amplified speech by listeners with hearing loss: acoustic correlates. *J Acoust Soc Am.* 1995;98:1388-99.
- Souza PE, Tremblay KL. Combining acoustic, electrophysiological and behavioral measures of hearing aids. Scottsdale, AZ: American Auditory Society; 2005.
- Kewley-Port D. Time-varying features as correlates of place of articulation in stop consonants. *J Acoust Soc Am.* 1983;73:322-35.
- Souza PE, Tremblay KL, Davies-Venn E, Kalstein L. Explaining consonant errors using short-term audibility. Salt Lake City, UT: Am Acad Audiol; 2004.
- Kraus N, Nicol T. Brainstem origins for cortical 'what' and 'where' pathways in the auditory system. *Trends Neurosci.* 2005;28:176-81.
- Young ED, Sachs MB. Representation of steady-state vowels in the temporal aspects of the discharge patterns of populations of auditory-nerve fibers. *J Acoust Soc Am.* 1979;66:1381-403.
- Worden FG, Marsh JT. Frequency-following (microphonic-like) neural responses evoked by sound. *Electroencephalogr Clin Neurophysiol.* 1968;25:42-52.
- Russo N, Nicol T, Musacchia G, Kraus N. Brainstem responses to speech syllables. *Clin Neurophysiol.* 2004;115:2021-30.
- Russo NM, Nicol TG, Zecker SG, Hayes EA, Kraus N. Auditory training improves neural timing in the human brainstem. *Behav Brain Res.* 2005;156:95-103.
- Johnson KL, Nicol T, Zecker SG, Kraus N. Developmental plasticity in the human auditory brainstem. *J Neurosci.* 2008;28:4000-7.
- Chandrasekaran B, Kraus N. The scalp-recorded brainstem response to speech: neural origins and plasticity. *Psychophysiology.* 2010;47:236-46.
- Cunningham J, Nicol T, Zecker SG, Bradlow A, Kraus N. Neurobiologic responses to speech in noise in children with learning problems: deficits and strategies for improvement. *Clin Neurophysiol.* 2001;112:758-67.
- Parbery-Clark A, Skoe E, Lam C, Kraus N. Musician enhancement for speech-in-noise. *Ear Hear.* 2009;30:653-61.
- Hornickel J, Skoe E, Kraus N. Subcortical laterality of speech encoding. *Audiol Neurotol.* 2009;14:198-207.
- Hemanth N, Manjula P. Representation of speech syllables at auditory brainstem. *J Indian Speech Hear Assoc.* 2012;26:1-13.
- Pittman AL, Stelmachowicz PG. Hearing loss in children and adults: audiometric configuration, asymmetry, and progression. *Ear Hear.* 2003;24:198-205.
- Sairam VVS, Manjula P. Long term average speech spectrum in Kannada. University of Mysore; 2002.
- Skoe E, Nicol T, Kraus N. Phase coherence and the human brainstem response to consonant-vowel sounds. In: Association for research in otolaryngology symposium. 2008.
- Dillon H, Katsch R, Byrne D, Ching T, Keidser G, Brewer S. The NAL-NL1 prescription procedure for non-linear hearing aids. National Acoustics Laboratories Research and Development, Sydney: National Acoustics Laboratories; 1998.
- Iles MW. Speaker identification as a function of fundamental frequency and resonant frequencies. University of Florida; 1972.
- Horii Y. Fundamental frequency perturbation observed in sustained phonation. *J Speech Hear Res.* 1979;22:5-19.
- Abbott E. A laryngographic study of voice quality. University of London; 1976.
- Ashmore JF. A fast motile response in guinea-pig outer hair cells: the cellular basis of the cochlear amplifier. *J Physiol.* 1987;388:323-47.
- Knight RT. Distributed cortical network for visual attention. *J Cogn Neurosci.* 1997;9:75-91.
- Krishnan A. Human frequency-following responses: representation of steady-state synthetic vowels. *Hear Res.* 2002;166:192-201.