

Conventional BTE vs RIC (receiver in the canal) BTE: A Comparative Study on Perceptual and Acoustic analysis of Speech and Music

Joseph D.S.¹ & Sandeep M.²

Abstract

Digital hearing aids are the newest innovation in hearing aids, providing more accuracy and higher quality sound than any other form of hearing aid on the market. This study focussed on a new digital technology named, Receiver in the canal hearing aid (RIC, which is similar to conventional digital hearing aids with the exception that the receiver of the hearing aid is placed inside the ear canal of the user and thin electrical wires replace the acoustic tube of the conventional behind the ear (BTE) aid. This is designed for more natural sound quality. The present study compared the efficacy of RIC and Conventional BTE hearing aids in the processing of speech and music samples. The efficacy was evaluated by comparing the speech and music samples processed through RIC and conventional BTE with that of natural samples. The comparisons were made both on subjective and objective analysis. Results of perceptual experiment showed that, speech is more natural through the RIC hearing aid compared to conventional BTE hearing aid. Acoustic analysis of speech through formant analysis showed that, formant frequencies (F1, F2, & F3) of /a/, /i/ & /u/ stimuli were better represented in RIC compared to conventional BTE hearing aid. Perceptual and acoustic experiment showed that music stimuli processed through both the hearing aid was almost similar. In both the hearing aids representation of music was poor compared to speech.

Key words: Receiver in the canal, hearing aids, processing.

Hearing aids are the amplification devices used by individuals with hearing loss. In terms of technology, they can be classified into analog and digital hearing aids (Sandlin, 2000). Digital hearing aids convert sound waves using exact mathematical calculations, which produce an exact duplication of sound. As a result, the sound quality produced by digital hearing aids is significantly higher than the quality of sound produced by analogue hearing aids (Dillon, 2001). Digital hearing aids are reported to have better signal quality compared to analog hearing aids (Wood & Lutman, 2004) which supports a lesser signal distortion in digital hearing aids.

Conventional digital BTE aids have a small plastic case that fits behind the pinna (ear) and provides sound to the ear via air conduction of sound through a small length of tubing, or electrically with a wire and miniature speaker placed in the ear canal. The delivery of sound to the ear is usually through an earmold that is custom made, or other pliable fixture that contours to the individual's ear.

Receiver In the canal (RIC) hearing aids are similar to the BTE aid. There is however one crucial difference: The 'receiver' of the hearing aid is placed inside the ear canal of the user and thin electrical wires replace the acoustic tube of the BTE aid. There are some advantages with this approach: Firstly, the sound of the hearing aid is arguably smoother than

that of a traditional BTE hearing aid. With a traditional BTE hearing aid, the amplified signal is emitted by the receiver which is located within the body of the hearing aid (behind the ear). The amplified signal is then directed to the ear canal through an acoustic tube, which creates a peaky frequency response. With a RITE hearing aid, the receiver is right in the ear canal and the amplified output of the hearing aid does not need to be pushed through an acoustic tube to get there, and is therefore free of this distortion. Secondly, RITE hearing aids can typically be made with a very small part behind-the-ear and the wire connecting the hearing aid and the receiver is extremely inconspicuous. For the majority of people this is one of the most cosmetically acceptable hearing device types. Thirdly, RITE devices are suited to "open fit" technology so they can be fitted without plugging up the ear, offering relief from occlusion (Kuk, 2008).

The main difference between the conventional BTE and RIC BTE is that the acoustic output of hearing aid receiver directly (ie, without going through earhooks and earmolds), given to the ear canal. And also a bandwidth that is broader than the standard configuration should be achievable.

Another important advantage of a RIC hearing aid, is that the earhook/tubing resonance that is typically seen in a BTE at 1000 Hz, 3000 Hz, and 5000 Hz when used without an earhook or an earmold tubing (RIC), disappears. Rather, the resonance peaks are replaced by a much smaller resonance peak at around 2500 Hz. The frequency response of the thin-wire RIC hearing aid is broader

¹ e-mail: divyasarah@gmail.com, ² Lecturer in Audiology, AIISH; email: msandeepa@gmail.com

than that of the thin-tube hearing aid, and with the use of newer receiver designs, it is possible that some newer RIC hearing aids have a broader frequency bandwidth which is likely to provide a richer and clearer sound quality than a more restricted bandwidth (Kuk, 2008).

Historically, the primary concern for hearing aid design and fitting is optimization for speech inputs. However, increasingly other types of inputs are being investigated and this is certainly the case for music. Whether the hearing aid wearer is a musician or merely someone who likes to listen to music, the electronic and electro-acoustic parameters described can be optimized for music as well as for speech. That is, a hearing aid optimally set for music can be optimally set for speech, even though the converse is not necessarily true. Similarities and differences between speech and music as inputs to a hearing aid are described. Many of these lead to the specification of a set of optimal electro-acoustic characteristics. Parameters such as the peak input-limiting level, compression issues—both compression ratio and knee-points—and number of channels all can deleteriously affect music perception through hearing aids. Regardless of the existence of a music program, unless the various electro-acoustic parameters are available in a hearing aid, music fidelity will almost always be less than optimal (Frank, 1982).

While a great number of hearing aid users – musicians among them – report excellent results from their respective devices, an acceptable reception of musical sound remains elusive for great numbers of others, and the applications for music and for speech sometimes appear to be at odds with one another, even though some of the newest and most sophisticated hearing aids are designed to take care of both music and speech without separate programmes and without any user-operated controls. Some individuals find that their otherwise helpful hearing aids tend to block out some of the defining sounds they are able to hear without those devices – the ping of a triangle, the growl of the double basses, the burr on the trombones – leaving them with a listening experience that is unsatisfying and frustrating, or at best a matter of compromises to which one becomes resigned, and perhaps looking for alternative ways of listening to music.

The improvement in speech perception may result from the better audibility provided in the high frequencies by the digital compression hearing device in comparison with the conventional aids. However, a subsequent study found that increasing the high-frequency gain in the conventional aids did not produce equivalent perceptual benefits is music.

In general, enhancing music perception has proved to be a good deal of a challenge than speech. With regard to speech, we are concerned basically

with intelligibility and voice recognition, while music involves both broader and finer concerns. In conversation, we may simply forget we are wearing effectively fitted hearing aids, but with music many listeners find, even after fastidious adjustment, that there is at least a trace of something mechanical, an absence of some nuance that gives life to the sound of specific instruments. High notes – from the violin's E string or a piccolo or a triangle – prove elusive, and huge boosts at the top all too often produce only noise and imbalance while leaving those elusive sounds utterly unrepresented.

Thus review of literature shows that there could be difference between conventional and RIC the aid output of both BTEs in terms of speech quality. So there is a need to check which hearing aids gives better perception of speech. There are only a few studies that included music in their quality judgment task (Franks, 1982). Most hearing aid users were not satisfied with their hearing aid for music listening. So there is a need to check which BTE model gives better music perception.

The aims of the present study are to: (1) To compare the speech processed through conventional BTE versus RIC hearing aids, on perceptual and acoustic analysis by normal hearing individuals.

Method

The study was conducted in two phases: preparation of stimulus and analysis of speech and Music samples

Preparation of Stimulus

Phase 1 included stimulus generation for the experiment. This was carried out in two stages; one, programming of the hearing aid and two, recording of the stimuli. Speech and music were the target stimuli. Fifty phonetically balanced Kannada words, developed by Vandana (1998) were used to prepare the target speech samples. Music samples included samples of violin and Mrudangam. These two musical instruments were taken based on the premise that sound from a violin is predominantly mid and high frequency, while the sound from a Mrudangam is predominantly low frequency. Both the samples were of 30 seconds duration. These words and music samples were processed through the 2 target hearing aids (a conventional digital & a RIC hearing aid) to prepare the test material for the study.

Procedure

Step 1- Programming of the Hearing Aids: Prior to recording, both the hearing aids were programmed for a hypothetical moderately severe (50 dB) flat hearing loss in all audiometric frequencies. The hearing aid was connected to a computer with programming software NOAH through Hipro. After

the hearing thresholds were fed into the software (NOAH), the hearing aids were programmed based on the NAL –NL1 prescriptive procedure.

Step 2 -Recording of the Test Stimuli: To record the words processed through the hearing aids, the words were initially fed into a computer. The audio output of the computer was routed into a calibrated diagnostic audiometer. The words were then played at 40 dB HL through a sound field speaker. The speaker was placed at 45 ° azimuth. One of the two digital hearing aids was placed in the client's position at 1 meter distance.

The receiver of the hearing aid was connected to a 2cc coupler. The other end of the coupler was attached to a sound level meter (SLM). The SLM in turn was connected to another computer which received and recorded the stimuli processed through the hearing aids. All the stimuli were generated with sampling rate of 44,100 Hz and 16 bit resolution. The so recorded words were then normalized to maintain the overall amplitude constant across the 50 target words. The words were then stored in a computer as individual files.

The same procedure was followed for the second hearing aid and for the recording of music samples. As a result, there were 3 sets of samples for further experimentation: Original natural speech and music samples, Speech and music samples processed through RIC BTE hearing aid, Speech and music samples processed through conventional BTE hearing aid.

Analysis of Speech and Music Samples

The speech and music samples were analyzed both on subjective and objective analysis.

Subjective or Perceptual Analysis: The samples were perceptually analyzed by fifty sophisticated listeners, who were in the age range of 18 to 30 years. To be the listeners in the present experiment, they had to fulfill the following criteria. Air conduction thresholds within 15 dB HL at octave frequencies between 250 Hz and 8 kHz and also had good speech identification scores (>90%) in quiet. This was tested using a calibrated GSI-61 audiometer. Normal result on immittance evaluation was also ensured using a calibrated GSI-TympStar immittance meter.

All the listeners were native speakers of Kannada and were blindfolded to the purpose of the study. The samples were randomly played to the listeners using an audio deck in a sound treated room and were asked to independently rate. Perceptual analysis was in terms of two parameters. (1) Speech identification scores. (2) Quality judgment

Speech Identification Scores: Firstly, speech identification scores for the processed stimulus were obtained from each subject. Only 1 half list of Vandana's Speech identification test was used while obtaining Speech identification scores for the words processed through each hearing aid.

Quality Judgment: The quality of speech and music sample were rated using a five point rating scale. The five point of the scale were: 1- Poor, 2- Fair, 3- Good, 4- Very Good and 5- Excellent.

Acoustical or Objective Analysis: Acoustic analysis included Long term speech spectrum (LTASS) and formant analysis. To do this, recorded speech (Vandana's word lists) and music samples (violin & mridangam) were fed into a computer with Vaghmi software (version-4V1). This was carried out for the samples processed through both conventional BTE and RIC BTE hearing aids also.

Spectral measures like frequency range and the peak frequency, at E01 (energy between 0 & 1Khz), E15 (energy between 1 & 5Khz), E02 (energy between 0 & 2Khz), E28 (energy between 2 & 8Khz), E58 (energy between 5 & 8Khz), α , β , and γ values were taken down from the LTASS.

Formant analysis was done by using Praat software (version-3.4). Spectrograms of only those recorded words which consisted any one of the three vowels; /a/, /i/ and /u/ were analyzed for F1, F2 and F3 formants. Formant analysis was not done for music samples.

The perceptual and acoustical data thus obtained was tabulated. Mean and standard deviation of the data were obtained for all the parameters analyzed. For the perceptual analysis of samples, Friedman's test was applied to check whether there were any significant difference between sample 1, sample 2 and sample 3.

Repeated measures ANOVA was done to check whether there were any significant differences between the 3 samples on the acoustic analysis.

Results

Results of Perceptual or Subjective Analysis

Three speech and six music samples were used in the study. The details of these samples are provided in Table 1. All the fifty listeners obtained 100 % scores speech identification scores for words in all three conditions. Hence there was no further analysis done in the speech identification scores. The 3 speech samples were perceptually rated by 50 sophisticated listeners on a 5 point rating scale. The compiled rating of the three speech samples by the 50 listeners is given in Table 2.

Table 1. Details of the samples used in the study

Samples Number	Condition
Sample 1	50 words of Vandana's list, one violin sample & one mrudhangam sample in the original, unmodified condition.
Sample 2	50 words of Vandana's lists, one violin sample & one mrudhangam sample processed through RIC BTE hearing aid.
Sample 3	50 words of Vandana's lists, one violin sample & one mrudhangam sample processed through conventional BTE hearing aid

Table 2. Total number of listeners who assigned a particular rating for the 3 speech samples

Rating	No of listeners		
	Sample 1	Sample 2	Sample 3
5(Excellent)	26	-	-
4(Very Good)	24	28	12
3(Good)	-	20	25
2(Fair)	-	2	13
1(Poor)	-	-	-

As in Table 2, the sample 2 & 3 perceptually rated were lower than that of sample 1. Also, sample 3 was rated lower than sample 2. None of the samples were rated poor. Rating of speech samples were statistically compared on Friedman's test. Results of the test revealed that there is a significant difference ($p < 0.001$) across the rating of the 3 samples. However, it was important to obtain pair-wise comparison to verify the objectives of the study. Hence, Wilcoxon Signed Ranks test was used and the results showed a significant difference between all 3 pairs: samples 1 & 2, 1 & 3 and, also 2 & 3.

Perceptual Analysis of Music: The rating scale used for the perceptual rating of music was same as that of speech. Combined rating was given for the samples of two musical instruments (Mrudhangam & Violin). The compiled rating of the 3 Music samples is as given in Table 3.

The data showed that, overall music samples were rated lower than speech samples. Among the 3 music a sample, sample1 was rated the best followed by sample 2 and sample 3. Within sample 2 & 3 sample 2 was rated higher. Friedman's test was used

to statistically compare the rating of 3 music samples and the results showed a significant difference in the rating of the 3 samples. The pair-wise significant difference was tested on Wilcoxon Signed Ranks Test. The results showed that there is significant difference ($p < 0.001$) between samples 1 & 2 and 1 & 3. However, there was no significant difference between sample 2 & 3.

Table 3. Total number of listeners who assigned a particular rating* for the 3 samples

Rating	Number of listeners		
	Sample 1	Sample 2	Sample 3
5	13	-	-
4	32	-	-
3	5	29	25
2	-	21	25
1	-	-	-

Results of acoustic or objective analysis

Speech and Music samples were also subjected to acoustic analysis. The acoustic analysis was done by using Long term average speech spectrum (LTASS) and formant analysis.

LTASS of Phonetically Balanced Words: LTASS was done to measure the energy distribution across different frequency range. Energy in the frequency range between 0 to 1kHz (E01), 1 to 5 kHz (E15), 0 to 2kHz (E02), 2 to 8 kHz (E28), 5 to 8kHz (E58) were taken down from the LTASS. In addition to these measures, α (E01/E15), β (E02/E28), & γ (E01/E58) were determined.

In general, the data showed that the means of sample 2 & sample 3 were different from that of sample 1. Also, sample 3 was more deviant from sample 1 compared to sample 2. The mean and standard deviation of the measures of LTASS for the 50 phonetically balanced words are given in Table 4.

Repeated measures ANOVA was done to evaluate whether the LTASS measures were statistically different across the 3 samples. Results showed that there is significant difference in E01, E15, E28, E58 and α across the 3 speech samples. Results of repeated measures ANOVA for the measure of LTASS are given in Table 5. Mean measures were not statistically different in E02, β and γ . The pair-wise significant difference was tested on Bonferroni and the results are depicted in Table 6.

Table 4. Mean and standard deviation of LTASS for phonetically balanced words

LTASS Measures	Sample1 Mean (SD)	Sample2 Mean (SD)	Sample3 Mean (SD)
E01	84.64 (1.68)	79.81 (4.01)	77.80 (3.48)
E15	76.22 (5.09)	84.61 (4.40)	85.85 (4.08)
E02	82.47 (16.03)	85.65 (3.67)	86.77 (4.38)
E28	70.38 (4.17)	70.28 (4.35)	74.1 (5.10)
E58	57.32 (5.53)	50.07 (4.10)	46.42 (4.42)
α	8.7 (5.68)	-3.2 (4.54)	-7.03 (4.57)
β	14.88 (4.49)	11.88 (4.82)	20.00 (20.38)
γ	26.14 (5.97)	29.97 (4.26)	29.81 (4.48)

Table 5. Repeated measures ANOVA results for the measure of LTASS for the words

LTASS Measures	df(Error df)	F value	P
E01	2 (46.3)	33.45	0.000
E15	2 (44.8)	66.27	0.000
E28	2 (45.1)	10.85	0.000
E58	2 (41.7)	38.38	0.000
α	2 (42.7)	152.6	0.000

The results can be summarized as follows; There is a significant difference in E58 and α value between all 3 pairs: 1 & 2, 1 & 3 and 2 & 3 ($p<0.001$). There is a significant difference in E01 and E15 between the samples 1 & 2 and 1 & 3 ($p<0.001$) and no significant difference between sample 2 & 3 ($p>0.05$). There is a significant difference in E28 between the samples 1 & 2 and 2 & 3 ($p<0.001$) and, no significant difference between 1 & 3 ($p>0.05$).

Formant Analysis of Phonetically Balanced Words: Formant analysis was done by using Praat software. Formant analysis was carried out only for the speech samples. First formant (F1), second formant (F2) and third formant (F3) of vowels /a/, /i/

Table 6. Results of Bonferroni Post hoc for LTASS measures of speech

LTASS Measures	Samples	Sample 1	Sample 2	Sample 3
E01	1	NS	S	S
	2	S	NS	NS
	3	S	NS	NS
E15	1	NS	NS	NS
	2	S	S	S
	3	S	NS	NS
E02	1	NS	NS	NS
	2	NS	NS	NS
	3	NS	NS	NS
E28	1	NS	S	NS
	2	NS	NS	S
	3	NS	S	NS
E58	1	NS	S	S
	2	S	NS	S
	3	S	S	NS
α	1	NS	S	S
	2	S	NS	S
	3	S	S	NS
β	1	NS	NS	NS
	2	NS	NS	NS
	3	NS	NS	NS
γ	1	NS	NS	NS
	2	NS	NS	NS
	3	NS	NS	NS

* Note: S- $p<0.05$, NS – $p>0.05$

and /u/ were obtained from the spectrogram of respective words. Total no: of words with vowels /a/, /i/, and /u/ were 18, 8, and 12 respectively.

Formants of vowel /a/: The mean and standard deviation (SD) of F1, F2, & F3 of vowel /a/ across the 3 samples are shown in Table 7. In general, the mean data showed that all three formants (F1, F2, & F3) were different in sample 2 & sample 3 compared to that of sample 1. Furthermore, the mean frequencies in sample 3 were more deviant from sample 1 compared to sample 2.

Table 7. Mean and standard deviation (SD) of F1, F2& F3 of /a/ vowel

Sample	F1 in Hz Mean (SD)	F2 in Hz Mean (SD)	F3 in Hz Mean (SD)
1	895.94 (96.99)	1592.66 (176.28)	2714.33 (254.73)
2	981.72 (69.46)	1607.44 (173.22)	2732.38 (257.24)
3	1115.66 (46.85)	1913.50 (113.84)	2910.55 (375.58)

Repeated measures ANOVA was done to evaluate whether the observed mean differences in formant frequency across the 3 samples were statistically different. The results of ANOVA showed significant difference in F1 [(F (2, 31.9) = 49.486), p

<0.001)], and F2 [(F (2, 29.3) = 29.859), $p < 0.001$] across the 3 the speech samples. Statistically, there was no significant difference ($p > 0.05$) in F3 of vowel /a/ across the 3 samples.

The pair-wise significant difference was tested on Bonferoni and the results showed a significant difference in F1 of vowel /a/ between all 3 pairs: 1 & 2, 1 & 3 and, 2 & 3. Also, there was a significant difference on F2 of /a/ vowel between the samples 1 & 3 and 2 & 3, while sample 1 & 2 were similar. The results are schematically represented in Table 11.

Formants of vowel /i/: The mean and standard deviation (SD) of F1, F2 & F3 of vowel /i/ across the 3 samples is shown in Table 8.

In general, the mean data showed that all three formants (F1, F2, & F3) were different in sample 2 & sample 3 compared to that of sample 1. Also, the mean formant frequencies in sample 3 were more deviant from sample 1 compared to sample 2.

Repeated measures ANOVA was done to evaluate whether the observed mean differences in formant frequency across the 3 samples were statistically different. The results of ANOVA showed significant difference in F1 [(F (2, 12.5) = 15.170), $p < 0.001$] across the 3 the speech samples. Statistically, there was no significant difference in F2 & F3 of vowel /i/ across the 3 samples ($p > 0.05$). The pair-wise significant difference was tested on Bonferoni and the results showed a significant difference in F1 frequency of vowel /i/ between all 3 pairs: 1 & 2, 1 & 3 and 2 & 3.

Formants of vowel /u/: The mean and standard deviation (SD) of F1, F2& F3 of vowel /u/ across the 3 samples is shown in Table 9.

In general, the data showed that all three formants (F1, F2, & F3) were different in sample 2 & sample 3 compared to that of sample 1. Also, the formant frequencies in sample 3 were more deviant from sample 1 compared to sample 2. Repeated measures ANOVA was done to evaluate whether the observed mean differences in formant frequency across the 3 samples were statistically different.

The results of ANOVA showed significant difference in F1 [(F (2, 18.9) = 251.45), $p < 0.001$], F2 [(F (2, 16.7) = 46.56), $p < 0.001$] and F3 [(F (2, 13.9) = 19.42), $p < 0.001$] across the 3 the speech samples. The pair-wise significant difference was tested on Bonferoni and the results showed a significant difference in F1 of vowel /u/ between all 3 pairs: 1 & 2, 1 & 3 and 2 & 3. There is a significant difference on F2 and F3 of vowel /u/ between sample 1 & 3 and 2 & 3, and no significant difference between 1 & 2

samples ($p > 0.05$). This is schematically represented in Table 10.

Table 8. Mean and standard deviation of F1, F2 & F3 of vowel /i/

Samples	F1 in Hz Mean (SD)	F2 in Hz Mean (SD)	F3 in Hz Mean (SD)
Sample1	469.62 (47.82)	2773.50 (136.37)	3203.87 (254.73)
Sample2	782.87 (40.45)	2773.25 (154.41)	3187.87 (257.24)
Sample3	1016.62 (76.87)	2510.25 (352.64)	3362.12 (375.58)

Table 9. Mean and standard deviation of F1, F2 & F3 of/u/ vowel

Sample	F1 Hz Mean (SD)	F2 Hz Mean (SD)	F3 Hz Mean (SD)
1	550.00 (29.06)	1143.75 (303.81)	2841.33 (92.47)
2	677.83 (76.71)	1180.5 (190.82)	2821.66 (65.45)
3	1005.00 (51.03)	2019.41(323.39)	3039.58 (147.95)

Acoustic Analysis of Music samples: Only LTASS and no formant analysis were done in the acoustic analysis of music samples. In general, the data showed that LTASS measures of sample 2 & sample 3 were different from that of sample 1. The sample 2 & 3 were almost similar on LTASS measures. The LTASS measures for the music samples are given in Table 11.

Discussion

Perceptual/ Subjective Analysis for speech:

Perceptual analysis of speech was in terms of two parameters, speech identification scores and quality judgment. Speech identification scores were 100% for all the 50 listeners with all three samples, since they are normal hearing listeners. There are two possible reasons for this, One, both the hearing aid used in the present study are digital hearing aids.

Second, the subject who participated as listeners were normally hearing individuals. So even if there was distortion, they would have compensated for the missing information through auditory closure. In the quality rating of speech, among the three samples (sample1, sample2, sample3) Sample 1 was rated higher followed by sample 2 and sample 3.

In the present study, sample 2 was RIC BTE and sample 3 was conventional BTE. So it can be interpreted from the results that RIC BTE is better than conventional BTE in terms of the quality of sound. In other words distortion was more in conventional BTE compared to RIC hearing aids. However, these distortions were not significant to

affect the speech intelligibility. It might be due to the following reasons; the frequency response of the thin-wire RIC hearing aid is broader than that of the conventional BTE hearing aid (Kuk & Baekgaard, 2008). A broader bandwidth would invariably provide a richer and clearer sound quality than a more restricted bandwidth. Perceptions of fuller sound and greater clarity can be expected with broader bandwidths (Ricketts, Dittbner, & Johnson, 2008). In RIC hearing aids, sound of the hearing aid is arguably smoother than that of a traditional BTE hearing aid, the receiver is right in the ear canal and the amplified output of the hearing aid does not need to be pushed through an acoustic tube to get there, and is therefore free of this distortion and RIC

Table 10. Boneferroni post hoc results of F1, F2 & F3 of /a/, /i/ & /u/ vowels

Vowel	Formants	Sample	Sample 1	Sample 2	Sample 3
/a/	F1	Sample 1	N S	S	S
		Sample 2	S	N S	S
		Sample 3	S	S	N S
	F2	Sample 1	N S	N S	S
		Sample 2	N S	N S	S
		Sample 3	S	S	N S
	F3	Sample 1	N S	N S	N S
		Sample 2	N S	N S	N S
		Sample 3	N S	N S	N S
/i/	F1	Sample 1	N S	S	S
		Sample 2	S	N S	S
		Sample 3	S	S	
	F2	Sample 1	N S	N S	N S
		Sample 2	N S	N S	N S
		Sample 3	N S	N S	N S
	F3	Sample 1	N S	N S	N S
		Sample 2	N S	N S	N S
		Sample 3	N S	N S	N S
/u/	F1	Sample 1	N S	S	S
		Sample 2	S	N S	S
		Sample 3	S	S	N S
	F2	Sample 1	N S	N S	S
		Sample 2	N S	N S	S
		Sample 3	S	S	N S
	F3	Sample 1	N S	N S	S
		Sample 2	N S	N S	S
		Sample 3	S	S	N S

*Note: S- $p < 0.05$, NS- $p > 0.05$

Table 11. Results of LTASS measures for music stimuli

LTASS measures	Sample 1		Sample 2		Sample 3	
	Violin	Mrudhangam	Violin	Mrudhangam	Violin	Mrudhangam
E01	88.43	85.84	69.41	59.41	62.81	59.00
E15	88.66	72.41	76.10	64.23	79.96	64.61
E02	95.61	86.04	75.19	65.23	78.27	64.87
E28	74.37	61.05	62.84	51.67	55.42	46.58
E58	42.28	52.94	35.78	28.88	34.31	32.88
α	0.228	13.43	-7.17	-2.82	-9.38	-5.163
β	17.18	24.57	13.15	13.30	20.0	17.98
γ	36.16	32.90	30.14	30.52	28.49	21.573

hearing aids have a smoother and wider frequency response (Hallenbeck & Groth, 2008). Since RIC BTE is in-cooperated with open fit technology, the unoccluded ear canal is reported to retain its natural resonance characteristics, enhancing the response in the 2-3 kHz region and further enhancing sound quality (Mueller & Ricketts, 2006). Where as in conventional BTE the ear hook and ear mould affects the quality of the amplified signal.

Acoustic / Objective Analysis for speech: The acoustic analysis was done by using Long term average speech spectrum (LTASS) and formant analysis. In LTASS overall data showed that the means of sample 2 & sample 3 were different from that of sample 1. Also, sample 3 was more deviant from sample 1 compared to sample 2. This shows that the two hearing aids distorted the spectral properties of the signal. Also the distortion in the spectrum was more with conventional BTE than RIC BTE hearing aids. These distortions could have been the reason for reduced quality in these hearing aids as evident in the quality rating. However, these distortions are not to an extent to reduce the speech identification in a normal hearing individual.

Results of formant analysis also showed similar trend, that all three formants (F1, F2, & F3) were different in sample 2 & sample 3 compared to that of sample 1. Result showed, formant frequencies were higher through hearing aids compared to the original signal. Furthermore, the formant frequencies in sample 3 (conventional BTE) were more deviant from sample 1 compared to sample 2 (RIC BTE). The present study shows that first (F1), second (F2) and third (F3) formants of vowel /a/, /i/, & /u/ was represented better in receiver in the canal hearing aids compared to conventional BTE hearing aids, this might be due to the free of distortion created by the replacement of ear hook and ear mould by a thin wire and ear tip in the RIC hearing aid and also since the receiver is placed right at the ear canal, the transformed acoustic energy is directly delivered to the ear canal. Where as in conventional BTE hearing aids due to the ear hook and ear mould, there will be acoustic modification of the input signal. These changes are in terms of a peaky frequency or resonance at 1 kHz, 3 kHz and 5 kHz (Dillon, 2001). These peaks and trough especially the peaks in the gain frequency response could adversely affect speech intelligibility and quality of the amplified sound (Dillon, 2001), particularly in a hearing impaired individual.

Results of the integration of both the perceptual and acoustic analysis showed that processing of speech samples is better represented through the receiver in the canal hearing aid compared to conventional BTE hearing aid. This is in agreement

with the earlier studies (Kuk & Baekgaard 2008; Chalupper, & Kasanmascheff, 2008).

Perceptual/Subjective Analysis of Music: Studies shows that new RIC technology provides better and natural sound quality. Since there are many difference between speech and music in terms spectral and temporal characteristics. So there is a need to check whether the processing of music stimuli through the new RIC technology is similar to speech or better/ poor.

Overall music samples were rated lower than speech samples, because most of the hearing aids are operating based on speech acoustics. So the quality music perception through the hearing aids may be lower compared to speech stimuli. Among the 3 music samples, sample 1 was rated the best followed by sample 2 and sample 3. Most of the subjects rated Sample 2 & 3 almost similar.

In both the hearing aids speech is represented well compared to music, since there are many differences between speech and music stimuli, the representation of speech and music through the hearing aids might be different. These results are supported by the following studies. Speech and music differs in terms of four major factors, they are (1) the long-term spectrum (2) overall intensities, (3) crest factors, and (4) phonetic vs. phonemic perceptual requirements of different musicians (Chasin 2003; 2006).

Wolfe (2002) pointed out several differences between speech and music in terms of fundamental frequency, formant structure, temporal structure, silence and transient spectral details. Hearing aid producing the best speech intelligibility may not provide the best sound quality for music (Gabreilsson & Sjogren, 1979; 1982).

A hearing aid ideal for music perception can be programmed to have good speech intelligibility but the vice versa is not true. It is important to understand the programming and internal algorithm changes necessary for optimal listening to music with hearing aids. This requires the knowledge of the differences between speech and music (Chasin, 2003).

Mishra, Kunnathur, and Rajalakshmi (2004) they obtained significantly poorer scores from music processed through hearing aids. They emphasized that music programs available in the commercial hearing aids cannot improve the fidelity for music because they operate based on speech acoustics.

Acoustical/Objective Analysis of Music: Results of the measures of LTASS showed that sample 2 and sample 3 were almost similar i.e, the music samples

processed through receiver in the canal hearing aid and conventional BTE were almost similar.

With the integration of both perceptual and acoustic experiment, overall there was an agreement between the perceptual and acoustical analysis. And in the present study, the overall results showed that the speech was better processed through receiver in the canal hearing aid compared to conventional hearing aid, whereas representation of music stimuli was almost similar in both hearing aids.

Conclusions

Receiver in the canal hearing aid is better in terms of quality compared to conventional BTE hearing aids. Speech stimuli are represented well through Receiver in the canal hearing aid than conventional BTE hearing aids. In both the hearing aids music stimuli was represented poorly compared to speech stimuli.

References

- American National standard Institute. (1991). American National Standard maximum permissible ambient noise levels in audiometric test rooms. ANSI, S3.1, New York. American National Standard Institute.
- Chasin, M. (2003). Music and hearing aids. *Hearing Journal*, 56(7), 36-41.
- Chasin, M., & Russo, F.A. (2004). Hearing aids and music. *Trends in amplification* 8(4), 35-47.
- Chasin, M. (2006). Can your hearing aid handle loud music? A quick test will tell you. *The Hearing Journal*, 59(12), 22-24.
- Dillon, H. (2001). Hearing Aids Systems. In Dillon H. (ed). *Hearing Aids*. (48-73). New York, Thieme, Boomerang Press.
- Frank, J.R. (1982). Judgement of hearing aid processed music. *Ear and Hearing*, 23(1), 18-23
- Gabrielsson, A., & Sjogren, H. (1979). Perceived sound quality of hearing aids. *Scandinavian audiology*, 8, 159-169.
- Hallenbeck, S. A., & Groth. (2008). Thin-tube and receiver-in-canal devices: There is positive feedback on both!! *The Hearing Journal*. 61(1), 28-34.
- Kuk, F., & Baekgaard. (2008). Hearing aid selection and BTEs: choosing among various "open-ear" and "receiver-in-the-canal" options. *Hearing Review*. 15(3), 22-36.
- Mirsha, S.K., Kunnathur, A., & Rajalakshmi, K. (2004). Why hearing aid and music seem to mix like oil and water? *National Symposium on Acoustics*, Mysore.
- Sandlin, R. E., (2000). *Textbook of Hearing Aid Amplification*, San Diego: Singular public group, Inc.
- Sjolander, Holmberg (2009). Broader Bandwidth Improves Sound Quality for Hearing-Impaired Listeners, *Hearing Review*.
- Vandana S. (1998). *Speech identification test For Kannada speaking children*. Unpublished Independent project. University of Mysore. India.
- Wood, S. A., & Lutman, M. E. (2004). Relative benefits of linear analogue and advanced digital hearing aids. *International Journal of Audiology*, 43(3), 144-155.
- Wolfe, J. (2002). Speech and music, acoustics and coding, and what music might be 'for'. Paper presented at *The 7th International Conference on Music Perception and Cognition*, Sydney, Australia.