

Recognition of Oriya Monosyllabic PB Words in Presence of Temporally Variant Noise

Biswajit Sadangi & Asha Yathiraj*

Abstract

The study attempted to determine the effect of temporally variant noise on recognition of monosyllabic PB words. Thirty Oriya speaking individuals with normal hearing were evaluated. Their speech identification scores were determined in the presence of speech shaped noise. The different conditions included evaluation of speech identification in the presence of continuous noise, 16 Hz interrupted noise and, 32 Hz modulated noise at two different SNRs (0 dB and -5 dB). In addition, their speech identification in a quiet condition was also obtained. No significant difference was found between the speech identification scores in the continuous noise and 32 Hz interrupted noise condition at 0 dB SNR, whereas with the 16 Hz interrupted condition a significant difference was seen.

Introduction

It has been proven time and again that the presence of background noise adversely affects the perception of speech. Kryter (1970) reported that when a speech signal is masked, either partially or completely by a burst of noise, its intelligibility changes in a complex manner. In research studies, various types of noises have been utilized to determine their influence on the intelligibility of speech. These noises can be categorized based on their temporal pattern or their frequency characteristics as well.

Assessing speech intelligibility in interrupted noise has been reported to reveal the auditory system's temporal ability to resolve speech fragments or get 'glimpses' or 'looks' of speech between the gaps of noise and to patch the information together to identify the specific speech stimuli (Miller, 1947).

Miller and Licklider (1950) reported about the intelligibility of monosyllables as a function of the interruption rate of noise. For interruption rates below 200 Hz, the speech intelligibility increased as the rate was lowered. The maximum intelligibility was reached at about ten interruptions per second. For low rates, the speech intelligibility dropped again because complete words were masked.

It was observed by Pollack (1955) that speech intelligibility decreased as the inter-burst level was increased for an interrupted masking noise of constant burst level. He found large improvements in speech intelligibility at lower repetition rates. Similarly, earlier Miller and Licklider (1950) observed that when the rate of interruption was 4 Hz or less, there was some loss of information because entire syllables and words were eliminated from the stimulus. However, once the interruption rate reaches 8 to 10 Hz, the words became as intelligible as un-interrupted speech.

* Professor of Audiology, AIISH, Mysore 570 006.

In contrast, it has been found by Dirks and Bower (1970), Dirks, Wilson and Bower (1969) and Miller and Licklider (1950) that word intelligibility did not change significantly when continuous speech was intermittently masked by white noise and interruption rates were varied from 1 to 100 Hz. They carried out the study using a 0 dB SNR.

It has been reported by Carhart, Tillman and Greetis (1969) that when the masker was the speech from a single talker or noise modulated periodically, speech intelligibility improved compared to when un-modulated noise was used even if the modulated and un-modulated noise had equal average energy (Dirks, Wilson and Bower, 1969; Festen and Plomp, 1990).

Results from the study of Smits and Houtgast (2007) revealed that individuals with normal hearing benefited from interruptions in noise while listening to digits in noise. The masking release was reported to be higher for the 16 Hz interruption than for 32 Hz interruption. The highest digit identification scores were obtained for 16 Hz modulated noise and lowest were for continuous noise.

The masking ability of a noise has also been noted to dependent upon the relation between the intensity of the speech and noise (Fant, 1960, cited in Silman and Silverman, 1991). For satisfactory communication, the signal-to-noise ratio was estimated to be +6 dB. When the criterion was not met, speech perception was noted to drop drastically. Moore (1996) found that at a 0 dB signal-to-noise ratio, word articulation scores reached 50%.

Groen (1969) evaluated the phoneme scores for individuals with hearing impaired and normal hearing individuals in presence of three SNRs (-5, 0 and +5 dB SNR). He observed that at -5 dB SNR the scores reduced drastically whereas at +10 dB the scores were significantly higher compared to at 0 dB SNR for hearing impaired individuals. For normal hearing group he found significantly lower scores at -5 dB SNR whereas no significant difference in scores were found for 0 dB and +10 dB. On the similar lines Kamlesh (1998) studied the effect of three SNR conditions (-5, 0 and +5 dB) on hearing impaired individuals using paired word (Kannada) recognition and questions. She reported that at adverse SNR the scores were significantly lower and with increase in SNR the scores improved. In the present study the similar results were obtained with the lowest scores for -5 dB SNR and better scores for 0 dB SNR in all the masking noise conditions. It is well documented in literatures about the need of speech perception tests in presence of noise for hearing aid selection (Carhart, 1965, cited in Plomp (1994); Tillman, Carhart and Olsen, 1970; Miller, Heise and Lichten, 1951; Stuart and Phillips, 1996).

From the literature, it is evident that speech perception of normal hearing individuals varies under fluctuating as well as steady-state noise. However, there is no consensus regarding its effect. There is a need to know if the speech identification

abilities of normal hearing individuals vary depending on whether the noise is interrupted temporally with different interruption rates. Further, noise in the environment occurs at different signal-to-noise ratios and often is fluctuating and not continuous. Hence, there is a need to know how individuals would perceive modulated noise in different signal-to-noise ratios. This would provide information about how individuals perform in a real life situation. Thus, the study aimed at comparing speech identification in the presence of continuous and temporally modulated speech shaped noise at two different signal-to-noise ratios (0 dB & -5 dB). The effect of two maskers having different temporal modulations was studied.

Method

Thirty individuals in the age range of 20 years to 50 years were evaluated in the study. All the individuals knew the dialect of Oriya spoken in Bhubaneswar region of Orissa. The participants did not have a history of any ear disorders. As they needed to provide an oral response, it was ensured that they did not have any speech or language disorders. In addition, in order to be included into the study, the participants had to have pure-tone thresholds within 20 dB HL. The participants whose speech identification scores were above 80% using the Oriya monosyllable PB wordlist, developed by Behera and Yathiraj (2007) were selected. All of the participants had an A' type tympanogram with reflexes present in both the ear and passed the Screening Checklist for Auditory Processing (SCAP) developed by Yathiraj and Mascarehans (2002).

A calibrated double channel, diagnostic audiometer, Orbiter 922 with TDH-39 headphone was used for the pure-tone air conduction testing and speech audiometry. A Radio Ear B-71 vibrator was used for estimating bone conduction thresholds. A calibrated middle ear analyzer, (GSI-Tympstar) provided tympanometry and reflexometry information. The speech and noise stimuli were presented through a Pentium 4 computer. The signals from the computer were routed to the audiometer.

Monosyllabic Phonemically Balanced (PB) Words of Oriya language, developed by Behera and Yathiraj (2007) was used. The test contained four half lists having phonemically balanced words. Each half-list had 25 monosyllabic words. The words were recorded digitally by a female Oriya speaker, who was fluent in the dialect of Oriya spoken in Bhubaneswar region of Orissa. The recordings were done using a Philips unidirectional microphone, connected to a Pentium IV computer, using the Adobe Audition 2.0 software. The recorded materials were scaled so that all the words were equally loud. Further, the four lists were randomized using a randomization table to form eight lists. Prior to each list a 1 kHz calibration tone was recorded. The materials were administered on 10 normal hearing Oriya speakers to ensure that the material was clearly recorded.

Speech shaped noise was generated with the Adobe Audition 2.0 software based on the parameters given by Silman and Silverman (1991). As suggested by them a broadband noise was filtered to have a frequency range of 250 Hz to 4000 Hz. The slope of the speech spectrum was +3 dB per octave from 250 Hz to 1 kHz and 12 dB per octave from 1 kHz to 4 kHz.

The speech shaped noise was further modulated to get 16 Hz and 32 Hz modulated noises. This was done using the MATLAB 7.0 software. The noises were then mixed with the recorded speech materials using the Adobe Audition 2.0 software. The following six conditions were thus generated: Continuous speech noise + speech at 0 dB SNR; 16 Hz modulated noise + speech at 0 dB SNR; 32 Hz modulated noise + speech at 0 dB SNR; Continuous speech noise + speech at -5 dB SNR; 16 Hz modulated noise + speech at -5 dB SNR and; 32 Hz modulated noise + speech at -5 dB SNR.

An example of a waveform for the word /kar/ is provided in Figure 1. Also given in Figure 1 are the waveforms for continuous noise, 16 Hz modulated noise, 32 Hz modulated noise, the combination of the test stimulus /kar/ with all the different types of noises used at 0 and -5 dB SNR.

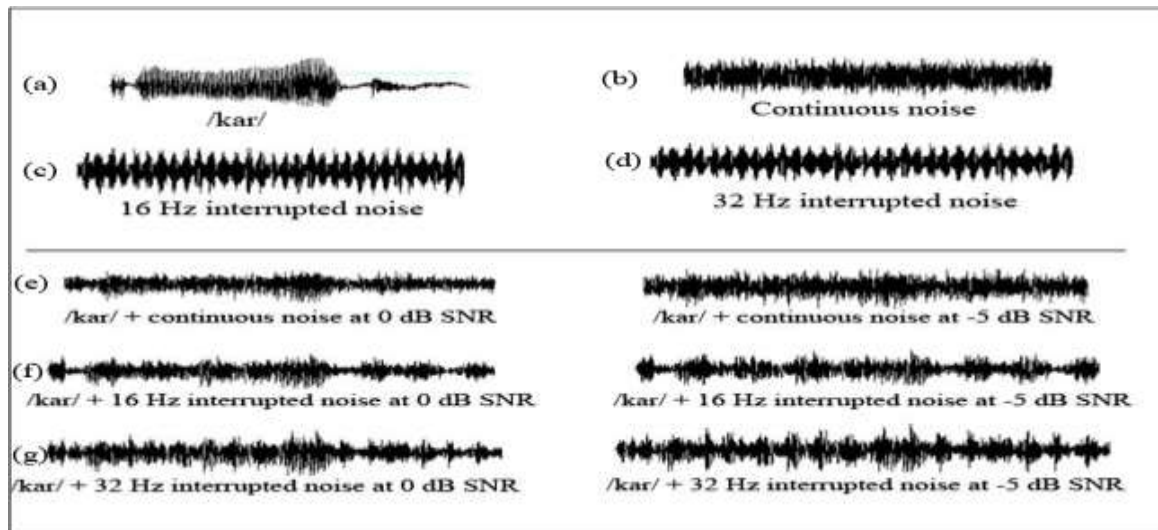


Figure 1: Waveforms of the (a) word /kar/, (b) continuous noise, (c) 16 Hz modulated noise (d) 32 Hz modulated noise; word /kar/ in combination with (e) continuous noise, (f) 16 Hz modulated noise and (g) 32 Hz modulated noise at 0 dB and -5 dB SNR

All the tests were carried out in a sound treated suite. The noise levels were within permissible levels as specified by ANSI S 3.1(1991).

The recorded speech materials were played on a Pentium-IV computer with the help of the Adobe Audition 2.0 software, routed through the audiometer and presented through headphones. All participants heard the speech signals at an intensity of 40 dB HL. In the two SNR conditions, the level of the signal was held constant, while the level

of the noise varied. Thus, in the -5 dB SNR condition the speech materials were presented at 40 dB HL and the noise was presented at 45 dB HL. The speech as well as noise was heard in the same ear. The choice of ear was randomized such that half the participants heard the signal through right ear and the other half through the left ear. The order in which each of the participants heard these lists were randomized to avoid any list effect. No participant heard the same list more than once.

The participants were instructed to repeat the words heard by them. The oral responses of the participants were scored. Every correct response was given a value of one and an incorrect response a score of zero. The data thus obtained on the thirty normal hearing participants were analyzed using the Statistical Package for Social Sciences (SPSS) software version 15. Repeated measure ANOVA and paired 't' test was carried out to determine the effect of the different masking conditions on speech perception.

The rationalized arcsine transform, developed by Studebaker (1985) was done to convert the speech identification scores into rationalized arcsine units (rau). This was done since it has been observed by Studebaker (1985) that speech identification scores are non-linear or additive. This was found to result in the critical difference between two speech identification scores being unequal. Hence, the available scores were converted to rau scores using the RATARC online rationalized arcsine transform program developed by Studebaker (1985). Thus, all statistical analyses were done for the word scores as well as for the rau scores.

Results and Discussion

The impacts of the following on the speech identification scores are discussed:

- Effect of different listening conditions (quiet, two SNR conditions and three masking conditions),
- Effect of signal-to-noise ratio (0 dB and -5 dB SNR) and,
- Comparison of quiet and masking conditions (continuous noise, 16 Hz and 32 Hz modulated noises at the two SNRs).

Effect of different listening conditions

The mean and the standard deviation of the speech identification scores in quiet and in the presence of noises were calculated separately. The mean speech identification scores were better for the quiet condition compared to the masking conditions. Among the difference noise conditions, better mean speech identification scores were obtained in the presence of 16 Hz modulated speech shaped noise at 0 dB SNR. In contrast, poorer scores were obtained for the 32 Hz modulated speech shaped noise at -5 dB SNR

condition. The mean and the standard deviation of the raw scores for the different conditions are given in Table 1 and Figure 2.

Also provided are the mean and standard deviation (SD) of the rau scores in Table 1 and Figure 3. Similar results were obtained for rau scores in all the listening condition as mentioned for the raw scores.

Table 1: Mean and Standard deviation (SD) for the speech identification (raw and rau) scores in different listening conditions

Listening Conditions	SNR	Raw scores		rau scores	
		Mean [#]	Standard deviation (SD)	Mean	Standard deviation (SD)
Quiet	-	23.90	1.09	102.51	9.74
Continuous noise	0 dB	19.93	2.66	79.37	11.86
Continuous noise	-5 dB	18.33	1.54	71.86	6.24
16 Hz modulated speech noise	0 dB	20.67	1.24	81.90	5.80
16 Hz modulated speech noise	-5 dB	19.53	1.11	76.74	4.63
32 Hz modulated speech noise	0 dB	19.03	1.79	74.92	7.91
32 Hz modulated speech noise	-5 dB	18.23	1.70	71.42	7.33

Maximum score = 25

In order to check whether there was a significant difference between the different noise conditions and also the different SNR conditions, two-way repeated measure ANOVA was done (3 noise conditions $\times$ 2 SNRs). The ANOVA results showed a significant main effect for the different noise conditions [$F(2, 58) = 8.52, p < 0.01$] and different SNR conditions [$F(1, 29) = 52.35, p < 0.01$]. It also showed a significant interaction effect between the different noise and SNR conditions [$F(2, 58) = 3.68, p < 0.05$].

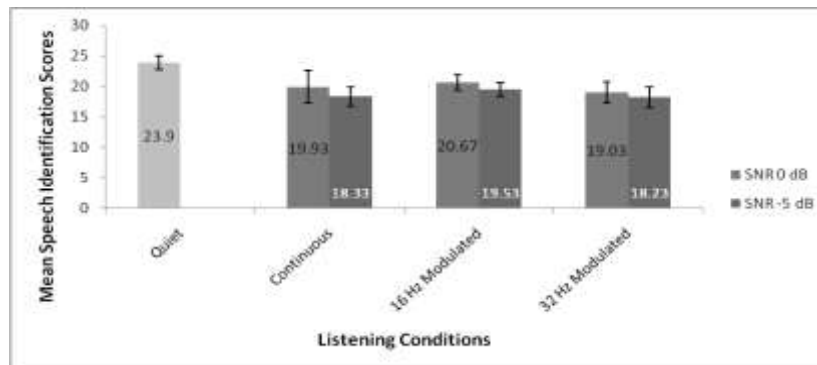


Figure 2: Mean and standard deviation for the speech identification raw scores across different listening situations and two different SNRs

ANOVA results for rau scores too showed a significant main effect for the different noise conditions [$F(2, 58) = 7.90, p < 0.01$] and different SNR conditions [$F(1, 29) = 53.45, p < 0.01$]. A significant interaction effect was also seen between the different noise and SNR conditions [$F(2, 58) = 4.39, p < 0.05$].

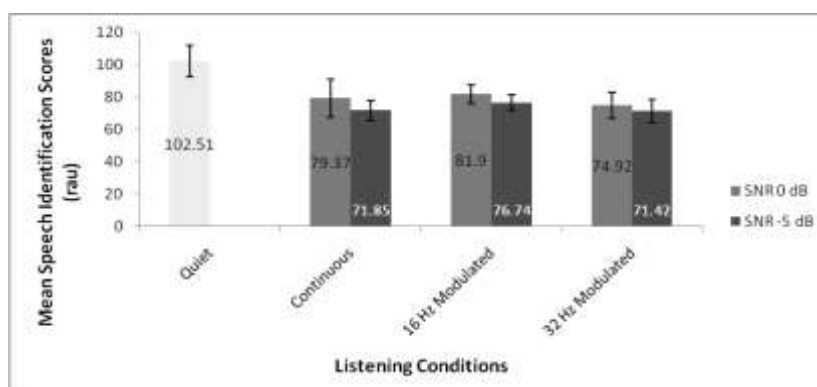


Figure 3: Mean and standard deviation for the speech identification rau scores across different listening situations and two different SNRs

Since the two SNR conditions were significantly different, separate one-way ANOVAs were done for the raw scores. This was done to see the significance of difference for different masking noise conditions at each of the two SNRs. A significant main effect was seen for the 0 dB SNR [$F(2, 58) = 6.99, p < 0.01$] and -5 dB SNR [$F(2, 58) = 9.24, p < 0.01$].

Similarly, on analysis of the rau scores a significant main effect was seen. This was observed for the 0 dB SNR [$F(2, 58) = 6.60, p < 0.01$] and -5 dB SNR [$F(2, 58) = 8.75, p < 0.01$] conditions.

Further, Bonferroni pairwise comparison was done to check the significance of difference between the different masked noise conditions. At the 0 dB SNR, a significant

difference between the 16 Hz modulated speech noise condition and 32 Hz modulated noise condition was observed. However, unlike the expected findings, there was no significant difference between the continuous and the modulated speech noise conditions (Table 2). On the other hand, the pairwise comparison for the -5 dB SNR revealed a significant difference ($p < 0.05$) between the continuous masking condition and the 16 Hz modulated noise condition. A similar significant difference was also observed between the two modulated conditions (16 Hz & 32 Hz) (Table 3). However, like that obtained at 0 dB SNR condition, no significant difference ($p > 0.05$) was seen for the continuous and the 32 Hz modulated noises.

Table 2: Pairwise comparison of different listening conditions at 0 dB SNR for raw scores

0 dB SNR		
Masking noise condition	16 Hz modulated	32 Hz modulated
Continuous	$p > 0.05$	$p > 0.05$
16 Hz modulated	-	$p < 0.05$

Table 3: Pairwise comparison of different listening conditions at -5 dB SNR for raw scores

-5 dB SNR		
Masking noise condition	16 Hz modulated	32 Hz modulated
Continuous	$p < 0.05$	$p > 0.05$
16 Hz modulated	-	$p < 0.05$

Similar results were found for the rau scores at both SNRs. Probably, identical results were obtained for the raw and rau scores since the values obtained in the present study did not contain scores along the entire range of scores. Most of the speech identification scores, across all conditions, were concentrated only in the upper extreme of the range.

It has been reported by Miller and Licklider (1950) and Gustafsson and Arlinger (1994) that at higher modulation rates, the release from masking was less. Hence, the speech intelligibility at higher modulation rates such as 32 Hz tended to be similar to be that observed for continuous or steady-state noise. In contrast, it has been reported that for maskers with lower modulation rates of 10 Hz (Miller and Licklider, 1950; Gustafsson and Arlinger, 1994) and 16 Hz (Smits and Houtgast, 2007), the release of masking was more, resulting in better speech perception compared to continuous noise maskers.

In the present study, such a release from masking was observed only at the -5 dB SNR and not at 0 dB SNR. This indicates that only at a more adverse SNR condition, did

the release of masking occur. Unlike the expected finding, no release in masking was obtained at the 0 dB SNR condition, even when the participants with extreme scores were eliminated. Hence, participant variability could not have accounted for the lack of release from masking at 0 dB SNR. Thus, it can be construed that release from masking is not solely dependent on the modulation rate but also on the SNR.

Further, in the present study the scores obtained using continuous masking noise were lower than that obtained with the 16 Hz modulated noise at -5 dB SNR. However, at the same SNR, no improvement was seen for the 32 Hz modulation rate compared to the continuous noise. This is in agreement with the results of various previous studies, where the masking release was reported to be higher for modulated noise than for continuous noise (Miller and Licklider, 1950; Gustafsson and Arlinger, 1994). It has also been reported by Smits and Houtgast (2007) that at higher modulation rates of 32 Hz, the noise functions similar to continuous noise and the advantage from release of masking does not occur.

Effect of signal-to-noise ratio (SNR)

From the mean values given in Table 1, it can also be observed that the speech identification scores were higher for the 0 dB SNR condition and lower for the -5 dB SNR condition. This was observed for continuous masking noise and both modulation masking (16 Hz and 32 Hz) conditions.

Paired sample 't' test was done to see if these differences in mean scores across the two SNRs were significantly different. The paired sample 't' test revealed a significant difference between the speech identification scores at the two different SNRs. This significant difference between 0 dB and -5 dB SNR ($p < 0.01$) was present in all the three masking conditions (continuous noise, 16 Hz modulated noise & 32 Hz modulated noise). This is evident from the information given in Table 4. Similar results were seen for the rau scores too.

Table 4: Significance of difference of different masking conditions at the two SNRs

Listening Conditions	SNR	Mean[#]	Standard deviation (SD)	't' value
Continuous noise	0 dB	19.93	2.66	4.94**
	-5 dB	18.33	1.54	
16 Hz modulated speech noise	0 dB	20.67	1.24	5.78**
	-5 dB	19.53	1.11	
32 Hz modulated speech noise	0 dB	19.03	1.79	5.17**
	-5 dB	18.23	1.70	

** Significant at $p < 0.01$

Maximum score = 25

The finding of the present study is in consonance with that of Groen (1969) and Kamlesh (1998). They too reported that at higher SNR the speech identification scores were higher and at a lower SNR of -5 dB the scores dropped. This drop in scores has been attributed to the masking that occurs at lower SNRs. Similar findings have also been noted by Olsen, Olofsson and Hagerman (2005) while using different other signal-to-noise ratio.

Effect of quiet Vs. masking conditions

The mean and standard values given in Table 1 clearly reveal that the speech identification scores were comparatively higher for the quiet condition than for any masking condition (continuous or modulated). Amongst the masking conditions, highest scores were obtained for the 16 Hz modulated noise at 0 dB SNR condition and the lowest scores were obtained for the 32 Hz modulated noise at -5 dB SNR condition.

To check for a significant difference between the quiet and the different masking conditions, paired 't' test was done. A significant difference between the quiet and all the different masking conditions was obtained. From Table 5 it can be noted that, the scores obtained in the quiet condition were significantly higher than that obtained with any of the masking conditions. Thus, irrespective of whether the masking noise had a modulation rate of 16 Hz / 32 Hz or had a SNR of 0 dB / -5 dB, it resulted in significantly lower scores than the quiet condition. Similar results were obtained with the rau scores as well.

Table 5: Significance of difference between the quiet and different noise conditions

Listening Conditions	Mean[#]	't' value
Quiet	23.90	7.67 ***
Continuous noise at 0 dB SNR	19.93	
Quiet	23.90	15.42 ***
Continuous noise at -5 dB SNR	18.33	
Quiet	23.90	15.60 ***
16 Hz modulated noise at 0 dB SNR	20.67	
Quiet	23.90	18.79 ***
16 Hz modulated noise at -5 dB SNR	19.53	
Quiet	23.90	13.13 ***
32 Hz modulated noise at 0 dB SNR	19.03	
Quiet	23.90	15.09 ***
32 Hz modulated noise at -5 dB SNR	18.23	

*** Significant at $p < 0.001$

Maximum score = 25

From the findings of the present study it can be noted that in the presence of masking noise, speech identification scores in normal hearing adults dropped drastically.

Even the least of the masking conditions (16 Hz at 0 dB SNR) was highly significantly different from the perception obtained in the quiet situation. The findings of the current study with regard to the comparison between the quiet and masking conditions are in agreement with that reported by Miller and Nicely (1955).

Though the present study was carried out with adult normal hearing participants, it can be construed that similar noise conditions would have a much more adverse affect on children. Studies in literature, comparing the performance of children with adults on speech intelligibility in noise, have shown that the former group performs poorer than the latter group (Elliott, 1979; Newman & Hochberg, 1983; Nitttrouer & Boothroyd, 1990). Further, noise levels in classrooms have been found to range from +35 dB to -10 dB as reported by Nebelek and Pickett (1974). It is possible that the intermittent noise present in the classrooms would have a highly negative effect on speech perception and hence the learning abilities of children. Thus, it is essential that noise levels should be much lower than what has been utilized in the present study in order to enable children to perceive speech effectively.

From the findings of the present study, it can also be extrapolated that if individuals with normal hearing are adversely affected with masking noise, those with a hearing loss are likely to be more adversely affected. Barrenas and Wikstrom (2000), Elpern (1960), Schneider and Daneman (2007), and Ross, Huntington, Newby and Dixon (1965) have reported that those with hearing impairment are likely to have more difficulty in speech perception in the presence of noise.

Conclusions

From the findings of the study it was observed that while a release from masking occurred at the more adverse SNR (-5 dB) no such release occurred at the lower SNR (0 dB). Thus, it can be inferred that release from masking is dependent on both modulation rate as well as SNR.

Further, it was observed that there was a significant difference between the two modulated masking noise condition (16 Hz and 32 Hz) at both at 0 dB and -5 dB SNR. Higher scores were obtained for the 16 Hz modulation condition. There was also a significant difference in the performances between the two SNRs (0 dB and -5 dB), with the scores being higher for the 0 dB SNR condition.

References

ANSI S3.1 (1991). Criteria for maximum permissible ambient noise during audiometric testing. NY: American National Standards Institute.

- Barrenas, M.L. & Wikstrom, L. (2000). The influence of hearing and age on speech recognition scores in noise in audiological patients and in the general population. *Ear and Hearing*. 21, 569-77.
- Behera, S. & Yathiraj, A. (2008). Development of monosyllabic speech materials in Oriya language. *Student research at AIISH Mysore articles based on dissertation done at AIISH, Vol : II (2003-2004) pp-41-52*.
- Carhart, R., Tillman, T. W. & Greetis, K. (1969). Binaural masking for speech by periodically modulated noise. *Journal of the American Society of Audiology*. 39, 1037-1050.
- Dirks, D.D. & Bower, D. (1970). Effect of forward and backward masking on speech intelligibility. *Journal of the American Society of Audiology*. 47, 1003-1008.
- Dirks, D.D., Wilson, R. H. & Bower, D. (1969). Effects of pulsed noise on selected speech materials. *Journal of the American Society of Audiology*. 46, 898-906.
- Elliott, L.L. (1979). Performance of children aged 9 to 17 years on a test of speech intelligibility in noise using sentence material with controlled word predictability. *Journal of the American Society of Audiology*. 66, 651-653.
- Elpern, B.S. (1960). Differences in difficulty among the CID W-22 auditory tests. *Laryngoscope*. 70, 1560-5.
- Festen, J.M. & Plomp, R. (1990). Effects of fluctuating noise and interfering speech on the speech-reception threshold for impaired and normal hearing. *Journal of the American Society of Audiology*, 88, 1725-1736.
- Groen, J. J. (1969). Social hearing handicap: its measurement by speech audiometry in noise. *International Audiology*, 8, 182-183.
- Gustafsson, H.A. & Arlinger, S.D. (1994). Masking of speech by amplitude-modulated noise. *Journal of the American Society of Audiology*. 95, 518-29.
- Kamalesh. (1998). Criteria for prescription of hearing aids. Unpublished master's dissertation, University of Mysore, India.
- Kryter, K.D. (1970). The effects of noise on man: Environmental sciences'. (2nd ed), pp: 234-35.
- Miller, G.A. (1947). Sensitivity to changes in the intensity of white noise and its relation to masking and loudness. *Journal of the American Society of America*, 19(4), 609-619.

- Miller, G.A. & Licklider, J.C.R. (1950). The intelligibility of interrupted speech. *Journal of the American Society of America*, Vol. 22(2), 167-173.
- Miller, G.A. & Nicely, P.E. (1955). An analysis of perceptual confusions among some English consonants. *Journal of the American Society of America*, 27, 338-352.
- Miller, G.A., Heise, G.A. & Lichten, W. (1951). The intelligibility of speech as a function of the context of the test materials. *Journal of Experimental Psychology*. Vol.41(5):329-35.
- Moore, B.J. (1996). Perceptual consequences of cochlear hearing loss and their implications for the design of hearing aids. *Ear and Hearing*.17.133-160.
- Nebelek, A.K. & Pickett, J.M. (1974). Reception of consonants in a classroom as affected by monaural and binaural listening, noise, reverberation, and hearing aids *Journal of the American Society of America*. Vol. 56(2):628-39.
- Newman, A.C. & Hochberg, I. (1983). Children's perception of speech in reverberation. *Journal of the American Society of America*. 73: 2145-2149.
- Nittrouer, S. & Boothroyd, A. (1990). Context effects in phoneme and word recognition by young children and older adults. *Journal of the American Society of America*. Vol. 87(6):2705-15.
- Olsen, H.L., Olofsson, A. & Hagerman, B. (2005). The effect of audibility, signal-to-noise ratio, and temporal speech cues on the benefit from fast-acting compression in modulated noise. *International journal of audiology*. 44 (7):421-33.
- Plomp, R. (1994). Comments on evaluating a speech-reception threshold model for hearing-impaired listeners. *Journal of the American Society of Audiology*, 96, 586-9.
- Pollack, I. (1955). Masking by a periodically interrupted noise. *Journal of the American Society of America*, 27(2), 353-356.
- Ross, M., Huntington, D., Newby, H. & Dixon, R. (1965). Speech discrimination of hearing impaired individuals in noise: its relationship to other audiometric parameters. *Journal of audiological research*. 5, 47-72.
- Schneider, B.A., Li, L. & Daneman, M. (2007). How competing speech interferes with speech comprehension in everyday listening situations. *Journal of American Academy of Audiology*, 18(7), 559-572.
- Silman, S. & Silverman, C.A. (1991) Stimuli commonly employed in audiological tests. *Auditory diagnosis principles and applications*, (pp:1-9). Academic press, Inc.

- Smits, C. & Houtgast, T. (2007). Recognition of digits in different types of noise by normal-hearing and hearing-impaired listeners. *International journal of Audiology*, 46(3):134-44.
- Stuart, A. & Phillips, D.P. (1996). Word recognition in continuous and interrupted noise by young normal-hearing, older normal-hearing and presbycusis listeners, *Ear and Hearing*, 17, 478-489.
- Studebaker, G.A (1985). A 'Rationalized' Arcsine Transform, *Journal of Speech and Hearing Research*, 28, 455-462.
- Tillman, T. Carhart, R. & Olsen, W. (1970). Hearing aid efficiency in a competing speech situation. *Journal of Speech and Hearing Research*, 13, 789-811.
- Yathiraj, A. & Mascarehans, K. (2002). Effect of auditory stimulation in central auditory processing in children with APD. ARF project carried out at department of audiology in All India Institute of Speech & Hearing, Mysore.