

To the Graduate Council:

I am submitting herewith a thesis written by Carole Johnson entitled "The Pitch of Frequency Specific Nonsense Syllables and Their Identification and Quality Assessment Through Octave Filter Bands." I have examined the final copy of this thesis for form and content and recommend that it be accepted in partial fulfillment of the requirements for the degree of Master of Arts, with a major in Audiology and Speech Pathology.


Carl W. Asp, Major Professor

We have read this thesis and
recommend its acceptance:




Accepted for the Council:


Vice Provost
and Dean of The Graduate School

STATEMENT OF PERMISSION TO USE

In presenting this thesis in partial fulfillment of the requirements for a Master's degree at The University of Tennessee, Knoxville, I agree that the Library shall make it available to borrowers under rules of the Library. Brief quotations from this thesis are allowable without special permission, provided that accurate acknowledgment of the source is made.

Permission for extensive quotation from or reproduction of this thesis may be granted by my major professor, or in his absence, by the Head of Interlibrary Services when, in the opinion of either, the proposed use of the material is for scholarly purposes. Any copying or use of the material in this thesis for financial gain shall not be allowed without my written permission.

Signature Casale Johnson
Date June 10, 1986

THE PITCH OF FREQUENCY SPECIFIC NONSENSE SYLLABLES
AND THEIR IDENTIFICATION AND QUALITY ASSESSMENT
THROUGH OCTAVE FILTER BANDS

A Thesis
Presented for the
Master of Arts
Degree
The University of Tennessee, Knoxville

Carole Johnson

June 1986

This thesis is dedicated to Caesar.

ACKNOWLEDGMENTS

The author wishes to thank:

- her committee members: Carl W. Asp, Ph.D. (Chairman), Bernard Silverstein, Ph.D., and Sol Adler, Ph.D.;
- her parents, Irma and Arthur Johnson, for their ever-present love, understanding, encouragement, emotional and financial support throughout her education;
- her best friends -- Nancy Horton and Stephanie Pirajnikoff -- who were always there to listen and care;
- her friends in Knoxville -- Ernie Thomas, Gail Mason, Yogi and Megumi Takata, Hsaio-Chuan Chen, Maria Serrano, Mary Jones, Cheryl Watson, Laura Hammond, Dawn Butler, etc. -- who made Knoxville a good place to be;
- Dr. Jeffrey L. Danhauer, who continues to inspire, guide, and encourage her in her academic endeavors;
- and Dr. Bernard Silverstein for serving as speaker in this study.

ABSTRACT

Each of the nonsense syllables, /mumu/, /vovo/, /lalə/, /keke/, and /sisi/ have a consonant and vowel which are homogeneous with regard to a particular pitch category -- low, low-middle, middle, middle-high, and high, respectively -- of the tonality model developed by Asp (1975) and Asp, Berry, and Bessell (1977).

The purposes of this study were to have twelve normal-hearing subjects -- 10 female, 2 male -- ages 24-40 years ($\bar{X}$ = 29.75): 1) judge the pitch of the five unfiltered frequency specific nonsense syllables -- within a paired-comparison paradigm; and 2) identify and assess the quality of the same syllables filtered through six one-octave filter bands with 20-26 dB per octave slopes and center frequencies: 250, 500, 1,000, 2,000, 4,000, and 8,000 Hz.

For the rating of pitch, results indicated: 1) the nonsense syllables were ranked from high to low pitch -- /sisi/, /keke/, /lalə/, /vovo/, and /mumu/; 2) statistically significant differences were found between all adjacently ranked nonsense syllables; 3) a significant difference existed between the five categories of the tonality model; and 4) the results agreed with the tonality model.

Results of the identification task indicated high scores (92-100%) for all filtered nonsense syllables except through the 250 octave filter band. The task was too easy. A linear averaging technique was used to obtain mean estimated optimal center frequencies for all nonsense syllables from the center frequencies of adjacent octave bands

producing the highest identification scores (96-100%). These estimated optimal center frequencies for the identification task were: 1) /vovo/ - 1,167 Hz; 2) /lɔlɔ/ - 750 Hz; 3) /keke/ - 3,750 Hz; and 4) /sisi/ - 3,750 Hz. An estimated optimal center frequency was not calculated for /mumu/ because the bands were not adjacent.

In the quality ratings, the same linear averaging technique was used. The estimated optimal center frequencies were: 1) /mumu/ - 2,000 Hz; 2) /vovo/ - 1,500 Hz; 3) /lɔlɔ/ - 1,167 Hz; 4) /keke/ - 2,000 Hz; and 5) /sisi/ - 4,000 Hz.

Generally the estimated optimal center frequencies for both tasks were in agreement with Guberina (1964; 1972) and Asp and Guberina (1982) for /lɔlɔ/, /keke/, and /sisi/, but not for /mumu/ or /vovo/. However, further inspection of the data revealed better agreement with Guberina (1964; 1972) and Asp and Guberina (1982) when optimal frequency regions were compared rather than estimated optimal center frequencies.

With regard to comparing the tonality of the unfiltered nonsense syllables to the filtered optimal octaves, /sisi/ and /keke/ were ranked highest on pitch and had the highest optimal frequency regions. However, this relationship was not clear for /mumu/, /vovo/, and /lɔlɔ/.

TABLE OF CONTENTS

CHAPTER	PAGE
I. INTRODUCTION	1
Purpose	4
Experimental Questions.	4
Definitions	5
II. REVIEW OF THE LITERATURE	8
Studies Involving the Pitch of Speech	8
Research Involving Optimal Octaves.	17
Use of Optimal Octaves in Speech Detection Testing. . .	21
Summary	23
III. METHODS.	25
Listeners	25
Speech Stimuli and Ordering Procedure	25
Recording Procedure for the Speech Stimuli.	27
Instrumentation	30
One-Octave Bands and Presentation Levels.	31
Listening Sessions.	35
Spectrographic Analysis	38
IV. RESULTS AND DISCUSSION	40
Post-Experiment Questionnaire and Reliability	40
Results and Discussion.	50
V. SUMMARY AND CONCLUSIONS.	81
Summary	81
Conclusions	84
BIBLIOGRAPHY	86
APPENDIXES	90
APPENDIX A. SUBJECT DEMOGRAPHIC INFORMATION	91
APPENDIX B. PRESENTATION ORDERS	94
APPENDIX C. FREQUENCY RESPONSES OF THE HEADSETS	98
APPENDIX D. CENTER AND CUT-OFF FREQUENCIES OF THE ONE- OCTAVE FILTER BANDS	101
APPENDIX E. PRESENTATION MODES FOR SPEECH STIMULI	103
APPENDIX F. EXPERIMENTAL PROTOCOL	109
APPENDIX G. OCTAVE BAND MEASURES OF AMBIENT NOISE IN SOUND-TREATED SUITE AND ALLOWABLE LEVELS FOR AUDIOMETRIC TEST BOOTHS	121
APPENDIX H. SPECTROGRAPHIC ANALYSES OF SPEECH STIMULI USED IN THIS STUDY.	124

CHAPTER	PAGE
APPENDIX I. MATRICES OF SUBJECTS' RESPONSES	160
VITA	170

LIST OF TABLES

TABLE	PAGE
I. Pitch Characteristics of the Five Principal Vowels as Determined by Early Speech Scientists	9
II. The Ranking of the Pre-Vocalic Consonants According to Pitch in a Paired-Comparison Study	12
III. Rank Order of the Pitch of 30 Homogeneous Words and Their Grouping into Three Pitch Categories.	16
IV. Optimal Octaves for British English Consonants and Vowels.	20
V. Logotomes Used in Verbo-Tonal Audiometry with Their Optimal Octaves and Comparable Pure Tone.	22
VI. Results of the Post-Experiment Questionnaire for 12 Listeners	41
VII. Subjects' Number, Number of Circular Triads, and Coefficient of Consistency.	46
VIII. Rankings of Pitch for the Five Unfiltered Homogeneous Nonsense Syllables and Times Rated Higher	51
IX. Percent of Identification for Five Nonsense Syllables Filtered Through Six One-Octave Filter Bands.	57
X. Means and Standard Deviations of the Quality Ratings of Five Homogeneous Nonsense Syllables Filtered Through Six One-Octave Bands.	65
XI. Adjacent t-Test Procedure for Obtaining Optimal Frequency Regions for the Nonsense Syllable /mumu/. . . .	71
XII. Adjacent t-Test Procedure for Obtaining Optimal Frequency Regions for the Nonsense Syllable /vovo/. . . .	72
XIII. Adjacent t-Test Procedure for Obtaining Optimal Frequency Regions for the Nonsense Syllable /lɔlɔ/. . . .	73
XIV. Adjacent t-Test Procedure for Obtaining Optimal Frequency Regions for the Nonsense Syllable /keke/. . . .	74

TABLE	PAGE
XV. Adjacent t-Test Procedure for Obtaining Optimal Frequency Regions for the Nonsense Syllable /sisi/	75
XVI. Estimated Mean Optimal Center Frequencies for the Five Homogeneous Nonsense Syllables (in Hz)	76
XVII. Optimal Frequency Regions for the Five Homogeneous Nonsense Syllables (in Hz).	79

LIST OF FIGURES

FIGURE	PAGE
1. Pitch (Tonality) Model Containing Phonemes Within Five Pitch Categories.	14
2. Block Diagram for Recording of Speech Stimuli	29
3. Block Diagram of Instruments for Experimental Sessions. . .	32
4. Frequency Responses of a Band-Pass Filter on a SUVAG Lingua at 1K Hz with a 1L Octave, One-Octave, 1/2-Octave, and 1/3-Octave Setting.	33
5. Block Diagram of Instruments for Spectrographic Analysis of Speech Stimuli	39
6. Matrix of Responses for Subject 1	43
7. Examples of Consistent and Inconsistent (Circular Triad) Ratings	45
8. Composite Matrix of all Subjects' Ratings of the Five Homogeneous Nonsense Syllables on Pitch	48
9. Bar Graph of Identification Scores for "No Attenuator Adjusted" (Cross-hatched) and "Relative Equal Loudness" Modes (Solid) for /mumu/ Filtered Through Six One-Octave Filter Bands.	58
10. Bar Graph of Identification Scores for "No Attenuator Adjusted" (Cross-hatched) and "Relative Equal Loudness" Modes (Solid) for /vovo/ Filtered Through Six One-Octave Filter Bands.	59
11. Bar Graph of Identification Scores for "No Attenuator Adjusted" (Cross-hatched) and "Relative Equal Loudness" Modes (Solid) for /lala/ Filtered Through Six One-Octave Filter Bands.	60
12. Bar Graph of Identification Scores for "No Attenuator Adjusted" (Cross-hatched) and "Relative Equal Loudness" Modes (Solid) for /keke/ Filtered Through Six One-Octave Filter Bands.	61

FIGURE

PAGE

13.	Bar Graph of Identification Scores for "No Attenuator Adjusted" (Cross-hatched) and "Relative Equal Loudness" Modes (Solid) for /sisi/ Filtered Through Six One-Octave Filter Bands.	62
14.	Bar Graph of Quality Ratings of /mumu/ Filterer Through Six One-Octave Filter Bands	66
15.	Bar Graph of Quality Ratings of /vovo/ Filtered Through Six One-Octave Filter Bands	67
16.	Bar Graph of Quality Ratings of /lɔlɔ/ Filtered Through Six One-Octave Filter Bands	68
17.	Bar Graph of Quality Ratings of /keke/ Filtered Through Six One-Octave Filter Bands	69
18.	Bar Graph of Quality Ratings of /sisi/ Filtered Through Six One-Octave Filter Bands	70

CHAPTER I

INTRODUCTION

Boring (1942) reported that as early as the 1790's, researchers conducted experiments whose results enabled them to rank the five principal vowels from low to high pitch /u, o, α , e, and i/ and assign them frequency values. Similarly, more recently, researchers investigated normal-hearing listeners' perceptions of the pitch of both whispered and voiced vowels (Harbold, 1958; McGlone and Manning, 1975). Not only were the results of these studies in agreement with the work of the early investigators as cited by Boring (1942), but the ranking of both whispered and voiced vowels on pitch were nearly identical indicating that the resonance of the pharyngeal/oral cavity may be more important for the pitch of phonemes than fundamental frequency (F_0).

Some investigators (Peterson, 1971; Peterson and Asp, 1972) had 23 normal-hearing listeners judge the relative pitch of 23 pre-vocalic English consonants paired with a neutral vowel / α /. Statistical analyses revealed that the consonants can be ranked according to low to high pitch.

Asp (1975) and Asp, Berry, and Bessell (1977) developed a pitch model that assigned vowels and consonants into low, low-middle, middle, middle-high, and high pitch categories. In using this pitch model, words are selected that have the vowel and consonant(s) from

the same or adjacent pitch category. For example, the word "see" has /i/, a high pitched vowel, and /s/, a high pitched consonant. The word "see" has phonemes that are homogeneous in pitch (tonality), and is a high tonality word. Thus, it is a frequency specific word.

Asp, Berry, and Bessell (1978) and Bessell (1979) selected 10 words from each of the low, middle, and high pitch categories of the pitch model. Then normal-hearing listeners judged the pitch of the 30 words within a paired-comparison paradigm. From the listeners' judgments, a statistical analysis revealed that the tonality words can be grouped into significantly different low, middle, and high pitch categories.

To minimize confusion between the pitch of fundamental frequency (F_0) and the pitch of the spectrum, Asp (1972) used the word "tonality" to refer to the pitch of the spectrum or "spectral pitch." Listeners can perceive the tonality difference between phonemes, e.g., /i/ is perceived higher than /f/, and /s/ higher than /m/. The tonality or spectral pitch differences allows the ordering of phonemes from low to high pitch. Optimal octaves are also ordered from low to high pitch, e.g., /u/ = 200-400 Hz; /a/ = 800-1,600 Hz; and /i/ = 3,200-6,400 Hz. Since tonality and optimal octaves are ordered in a similar manner, it appears that optimal octaves may be the physical correlate of tonality.

Guberina (1964;1972) found that filtering a phoneme through one one-octave filter band provided better perception than when filtering through another one-octave filter band. For example, /i/ was

was identified better through the 3,200-6,400 Hz band than the 800-1,600 Hz band. Thus, the 3,200-6,400 Hz band was identified as the optimal octave for the vowel /i/ (Guberina, 1972). Guberina reported optimal octaves for the phonemes of at least six languages taught at the University of Zagreb (Asp, 1972). British English is taught as a foreign language at the Institute of Phonetics at the University of Zagreb. Studies investigating the optimal octaves of vowels (McKenney, 1971; McKenney and Asp, 1972; Miner and Danhauer, 1977) are in agreement with the optimal octaves identified by Guberina (1964, 1972).

Roberge (1979) compiled the optimal octaves of French, English, German, Italian, and Spanish phonemes from the work of others. Similarly, optimal octaves were determined by Spletzer and Asp (1981), using 20-26 dB per octave slopes, whereas Silverstein (1984) used 11-14 dB per octave slopes. Guberina (1964, 1972) reported the optimal octaves for consonant and vowel nonsense syllables which are used as speech stimuli in Verbo-tonal Audiometry.

Asp (1972) stated that the optimal octaves of phonemes can be arranged on a frequency continuum from 200 to 12,800 Hz. He proposed that this physical continuum can be thought of as a perceptual continuum similar to the ordering of the tonality of unfiltered phonemes. Theoretically, in the perception of unfiltered phonemes, the brain uses the frequency region of the optimal octave for the identification of each phoneme (Asp, 1972). Asp (1972) stated that

the optimal octave of the phoneme is responsible for its "spectral" pitch or tonality. Consequently, there must be a relationship between the pitch of unfiltered phonemes and their corresponding optimal octaves. However, no study has yet to simultaneously investigate the pitch of unfiltered speech stimuli and their optimal octaves. Clearly, such a study is warranted.

I. Purpose

The purposes of this study were to have normal-hearing listeners:

- 1) judge the pitch of five unfiltered nonsense syllables -- /mumu/, /vovo/, /lola/, /keke/, and /sisi/ -- with the consonant and vowel homogeneously grouped according to the pitch model (Asp et al., 1977); and
- 2) identify and assess the quality of the same five homogeneous nonsense syllables, filtered through six one-octave filter bands with center frequencies of 250, 500, 1,000, 2,000, 4,000, and 8,000 Hz with 20-26 dB per octave slopes.

II. Experimental Questions

Ratings of the Five Unfiltered Homogeneous Nonsense Syllables on Pitch

1. Can the five unfiltered homogeneous nonsense syllables be ranked on pitch?
2. Is there a significant difference among the ranks of the five unfiltered homogeneous nonsense syllables on pitch?

Identification and Quality Assessment of the Five
Homogeneous Nonsense Syllables Each Filtered
Through Six One-Octave Filter Bands

Identification

1. Are there "optimal" one-octave filter bands for the perception of each of the five homogeneous nonsense syllables resulting from the normal-hearing listeners' identification through phonetic transcription of the five homogeneous nonsense syllables each filtered through six one-octave filter bands?

Quality Assessment

1. Are there "optimal" one-octave filter bands for the perception of each of the five homogeneous nonsense syllables resulting from the normal-hearing listeners' quality ratings of the five homogeneous nonsense syllables each filtered through six one-octave filter bands?

III. Definitions

The following terms were defined to clarify usage in this study:

Heterogeneous Grouping: A grouping of phonemes such that the phonemes are selected from different pitch categories according to the pitch model (Asp et al., 1977). For example, the nonsense syllable /su/ is heterogeneously grouped according to tonality since

it is composed of a consonant from the high pitch category /s/ and a vowel from the low pitch category /u/ (Asp, 1985).

Homogeneous Grouping: A grouping of phonemes such that the phonemes are selected from the same or adjacent pitch category according to the pitch model (Asp et al., 1977). For example, the nonsense syllable /si/ is homogeneously grouped according to tonality since it is composed of a consonant /s/ and /i/ which are both from the high pitch category (Asp, 1985).

Logotomes: A frequency specific phonic group consisting of a consonant and vowel (e.g. /mu/) repeated twice (e.g. /mumu/). The CV nonsense syllable is repeated twice to simulate the normal rhythm and linguistic redundancy of speech (Asp, 1985).

Optimal Octaves: The one-octave filter band producing the highest intelligibility score and providing for the enhanced perception of a phoneme from all other octave bands. The optimal octaves of phonemes are based on the perception of normal-hearing listeners (Asp, 1985).

Perceptual Continuum: The arrangement of the 40 phonemes of English on a continuum from low to high pitch or tonality for the purpose of analyzing perceptual segmental errors and planning strategies for their correction in therapy (Asp, 1985).

Tonality: The pitch of a phoneme, syllable, or word whereby /i/ is perceived as higher in pitch than /u/, /si/ higher than /mu/, and /sis/ (i.e. "cease") higher than /~~w~~od/ (i.e. "wood"). The pitch of speech stimuli appears to be related to the "spectral

pitch" of phonemes, and may be described as the timbre of the phonemes (Asp, 1985).

Transfer: The ability of a hearing-impaired person to identify phonemes through a limited frequency band (i.e. their optimal field of hearing) even though the spectral pitch of those phonemes appear to be outside of this limited frequency band. The ability to perceptually transfer is developed through auditory training. Transfer occurs at the level of the brain. Transfer can be tested with Verbo-tonal Audiometry. Transfer can also be used to explain the direction of errors along the perceptual continuum (Asp, 1985).

CHAPTER II

REVIEW OF THE LITERATURE

I. Studies Involving the Pitch of Speech

Studies Involving the Pitch of Vowels

Boring (1942) reported that as early as the 1790's, researchers conducted experiments whose results enabled them to rank the five principal vowels from low to high pitch /u, o, a, e, and i/ and assign them frequency values. The four basic methods of experimentation used in these earlier studies were: 1) artificial production of vowels using tuning forks; 2) identification of the resonant frequency of the oral cavity during vowel production using tuning forks; 3) comparison of the pitch of spoken vowels to pure tones; and 4) Fourier analysis of phonograph recordings of vowels. Table I illustrates the results of early speech scientists' investigations of the pitch characteristics of the five principal vowels (Boring, 1942).

Harbold (1958) compared the rank order of pitch for whispered and voiced vowels. Normal-hearing listeners chose the vowel which was "highest" in pitch for each pair of vowels within a paired-comparison paradigm. The results of this procedure resulted in a rank order of both whispered and voiced vowels according to pitch. The rank orders of whispered and voiced vowels had a positive .97 correlation. Both rank orders had the /i/ vowel ranked highest in

Table I. Pitch Characteristics of the Five Principal Vowels as Determined by Early Speech Scientists

Researcher	Year	Method	/u/	/o/	/a/	/e/	/i/
Willis	1829	Synth.	Low	525	590-1400	2360-4200	6300
Donders	1857	Reson.	350	295	490	1095	1400
Helmholtz	1863	Reson.	175	470	940	1880+175	2350+175
Koenig	1870	Reson.	235	470	940	1880	3760
Hermann	1895	Phonogr.	525-590	590-660	660-820	1970+2100	2350-3525
Miller	1916	Phonogr.	325	470	920-1240	1950+690	3100+300
Jaensch	1913	Synth.	250	500	1000	2000	-----
Koehler	1910	Comp.	262	525	1050	2100	4200
Modell and Rich	1915	Comp.	280	600	1140	1900	3750
Rich	1919	Comp.	262	525	1065	-----	-----

Synthesis = Synthetic creation of vowel sounds.

Resonance = Analysis of resonance of oral cavity set to speak each vowel.

Comparison = Matching a tone with the vowel character.

Phonograph = Analysis of spoken vowel.

pitch and the vowel /u/ the lowest. These findings suggest that the resonance of the pharyngeal/oral cavity may be more important for the pitch of phonemes than fundamental frequency (F_0).

Similarly, McGlone and Manning (1975) used four male speakers to record the vowels /i, I, æ , u, and ɔ /, both whispered and voiced, within the syllable /h__v/. Normal-hearing listeners chose the syllable which was "highest" in pitch within each pair of syllables in a paired-comparison paradigm. This resulted in the ranking, from highest to lowest pitch, of /i, I, æ , u, and ɔ / for the voiced vowels and /i, I, æ , ɔ , and u/ for the whispered vowels. The difference in ranking for the / ɔ / and the /u/ vowels may be due to the fundamental frequency of the voiced vowels. Spectrographic analyses for the voiced vowels indicated: 1) little agreement between fundamental frequency (F_0) or F_1 measurements and the rank order of the vowels on pitch; 2) little agreement between F_0 or F_1 measurements of the vowels reported in earlier studies (Peterson and Barney, 1952) and the rank orders of the vowels on pitch; but 3) strong agreement between measurements of F_2 of the vowels in this study and their rank order of pitch. For the whispered vowels, only F_2 of the vowels was measured and there was a strong agreement between these measurements and the rank order of the whispered vowels on pitch. The authors concluded that the pitch of either whispered or voiced vowels is due more to cavity resonance (i.e. F_2), rather than F_0 .

Studies Involving the Pitch of Consonants.

Syllables, and Words

Guberina (1964) first studied the pitch of speech sounds when he selected subjects with different auditory pathologies and assessed their speech perception abilities. He demonstrated that subjects with conductive pathologies more easily identified /t/ and /k/, rather than /m/, /b/, or /p/. However, Guberina found that subjects with sensori-neural pathologies more easily identified /m/, /b/, and /p/ than /t/ or /k/, but could not perceive /ʃ/ or /s/. Guberina speculated that since conductive pathologies allowed for the perception of high frequencies while attenuating the low frequencies and sensori-neural pathologies do the opposite, then /m, b, and p/ must be low frequency sounds and /t, k, ʃ, and s/ must be higher frequency sounds.

Peterson (1971) and Peterson and Asp (1972) had normal-hearing listeners judge the relative pitch of 23 pre-vocalic English consonants paired with a neutral vowel /a/ within a paired-comparison paradigm. Statistical analyses of the data revealed that the consonants can be ranked from low to high pitch. Table II illustrates the 23 pre-vocalic consonants ranked according to the responses of all 23 listeners and also to the ten best listeners in the study. In addition, the rankings of the 23 pre-vocalic consonants according to pitch resulted in significant differences in the grouping of the consonants into low, low-middle, middle, middle-high, and high pitch categories. Significant differences were also found between voiced and voiceless

Table II. The Ranking of the Pre-Vocalic Consonants According to Pitch in a Paired-Comparison Study

23 Listeners		10 Best Listeners	
Consonant		Consonant	
Low	n	Low	m
	m		g
	g		n
	ʒ		b
	dʒ		r
	r		d
	b		w
	d		l
	w		dʒ
	j		v
	i		j
	p		ʒ
	v		hw
	θ		h
	t		p
	z		k
	tʃ		θ
	h		f
	t		tʃ
	hw		z
	k		t
	ʃ		ʃ
High	s	High	s

Source: Peterson (1971) and Peterson and Asp (1972).

consonants; the voiceless consonants being higher in pitch. The nasals were found to be significantly lower in pitch than both the fricatives and affricates. Further, a significant correlation was found between the rank ordering of the consonants on pitch and Poole's developmental norms for phonemes (Poole, 1934) suggesting that lower pitched phonemes may develop earlier than higher pitched phonemes.

Asp (1975) and Asp, Berry, and Bessell (1977) developed a pitch model that assigned vowels and consonants into low, low-middle, middle, middle-high, and high pitch categories (see Figure 1). This model was based on research that identified the optimal octaves and the rank order of pitch of both consonants and vowels in addition to clinical experience with the hearing-impaired. In using this pitch model, the authors select homogeneous phonemes from the same or adjacent categories in formulating frequency specific words. For example, a consonant and a vowel can be combined from the low pitch category to form the word "moo" (i.e. /mu/) or from the high pitch category to form the word "she" (i.e. /ʃi/).

Asp, Berry, and Bessell (1978) and Bessell (1979) had normal-hearing listeners judge the relative pitch of 30 words -- consisting of homogeneous consonants and vowels within the low, middle, and high pitch categories of the pitch model. Statistical analyses revealed that the words can be ranked from high to low pitch and grouped into significantly different low, middle, and high pitch

	LOW	LOW-MIDDLE	MIDDLE	MIDDLE-HIGH	HIGH
CONSONANTS	M N	M N P	K F T ʁ ɖ ʒ	t ʒ ɐ	z ʃ θ ʈʂ
	b	g d			
	r	v			
VOWELS	u	o	a	e	i
		ʊ	ə	ɛ	ɪ
	ou	av	ɔ	aɪ eɪ	

Figure 1. Pitch (Tonality) Model Containing Phonemes Within Five Pitch Categories.

Source: Asp, Berry, and Bessell (1977).

categories. Table III illustrates the ranking of the 30 words according to pitch and their grouping into significantly different low, middle, and high pitch categories. For example, the word "cease" from the high pitch category was rated most consistently higher in pitch than all other words and the word "wood" was most consistently rated lower in pitch than all other words.

Asp and Johnson (1985) investigated the effects of both homogeneously and heterogeneously grouping phonemes according to the pitch model (Asp et al., 1977). This included pairing a low /m/, a middle /a/, and a high /s/ consonant each with a low /u/, a middle /a/, and a high /i/ vowel. This resulted in three homogeneous (i.e. /mu, la, and si/) and six heterogeneous (i.e. /ma, mi, lu, li, su, and sa/) nonsense syllables. Ten normal-hearing young adult listeners selected the syllable which was "highest" in pitch in each pair within a paired-comparison paradigm. Results indicated: 1) the nonsense syllables could be ranked from high to low pitch (i.e., /si, li, mi, sa, la, su, ma, lu, and mu/); 2) both consonants and vowels affected the listeners' judgments although the vowels appeared to have a greater effect; 3) the rank order of the nonsense syllables had a +.90 correlation with the ordered predicated by a mathematical perceptual model of tonality developed by Asp and Johnson (1985); and 4) the rank order of pitch of the nonsense syllables was in agreement with past research.

Table III. Rank Order of the Pitch of 30 Homogeneous Words and Their Grouping into Three Pitch Categories

Rank	Word	Pitch Category
1.	Cease	  High
2.	Seek	
3.	Seat	
4.	See	
5.	Teach	
6.	Cheap	
7.	Six	
8.	Each	
9.	Sheet	
10.	She	
11.	Tack	  Middle
12.	Tap	
13.	Have	
14.	Cat	
15.	Hat	
16.	Fat	
17.	Hot	
18.	Lot	
19.	Add	
20.	Rock	
21.	New	  Low
22.	Now	
23.	Move	
24.	Blue	
25.	Room	
26.	Bow	
27.	Row	
28.	No	
29.	Bone	
30.	Wood	

Source: Asp, Berry, and Bessell (1978) and Bessell (1979).

II. Research Involving Optimal Octaves

Chiba and Kajiyama (1958) filtered the vowels /u, o, a, e, and i/ through band-pass filters and found that when the vowels were filtered through: 1) a 0-300 Hz band-pass filter, the vowels were perceived as /u/; 2) a 300-600 Hz band-pass filter, the vowels were perceived as /o/; 3) a 900-1,600 Hz band-pass filter, the vowels were perceived as /a/; and 4) a 2,000-3,000 Hz, a 2,500-4,000 Hz, and a 3,000-5,000 Hz band-pass filter, the vowels were perceived as /i/. The filter bands associated with each particular vowel corresponds to one of the formant regions of the vowels in the unfiltered form. Consequently, these findings suggest that these formants may be the most important frequency regions for the identification of each vowel. Further, Chiba and Kajiyama (1958) describe the vowels /u, o, and a/ as single formant vowels with identified frequency regions of 350 Hz, 500 Hz, and 750 Hz, respectively. However, the authors describe the vowels /e/ and /i/ as two formant vowels with principal frequency regions of 2,800 and 2,700 Hz, respectively. Black (1968) stated that these results are in agreement with Guberina (1964).

Guberina (1964, 1972) analyzed normal-hearing listeners' perceptions of vowels which had been recorded, electroacoustically filtered through different one-octave filter bands, and played through to listeners' earphones. The listeners' task was to identify what vowel they heard. Guberina found that the perception of a vowel can change to another vowel depending on the one-octave filter band

that is used. For example, the vowel /i/ is perceived as /i/ when filtered through the one-octave band 3,200-6,400 Hz, but as /u/ when filtered through the 200-400 Hz octave band. The octave band producing the highest intelligibility score for a vowel was designated as the optimal octave for the vowel. The following optimal octaves were identified: 1) /u/, 200-400 Hz; 2) /o/, 400-800 Hz; 3) /ɑ/, 800-1,600 Hz; 4) /e/, 1,600-3,200 Hz; and 5) /i/, 3,200-6,400 Hz.

McKenney (1971) and McKenney and Asp (1972) provided additional data supporting the theory of optimal octaves in vowel perception. These authors investigated the effects of band-pass filtering on normal-hearing listeners' perceptions of the five English vowels /u, o, ɑ, e, and i/ spoken by an adult male speaker of General American English. The vowels were band-passed through 11 overlapping one-octave filter bands with 35 dB per octave slopes with center frequencies ranging from 100 to 4,000 Hz. The filtered vowels were presented to 15 normal-hearing listeners at 70 dB SPL. Listeners were instructed to listen to the vowel and identify the vowels that they heard. Results revealed that the optimal octaves for the vowels in this study were in agreement with Guberina (1964, 1972). The experimental procedure was then repeated again, except with a female speaker. The optimal octaves identified with this speaker for the /o/ and /e/ vowels were not in complete agreement with Guberina (1964, 1972). McKenney (1971) and McKenney and Asp (1972) stated that this finding did not coincide with Guberina's theory that optimal octaves are

independent of speaker and fundamental frequency. Future research is needed to investigate the nature of this discrepancy.

More recently, Miner and Danhauer (1977) investigated the relation between formant frequencies and the theory of optimal octaves of vowel perception. The vowels /u, a, and i/ were produced by an adult male speaker of General American English, filtered through eight one-octave filter bands (80-160, 160-315, 315-630, 630-1,250, 1,250-2,500, 2,500-5,000, 5,000-10,000 and 10,000-20,000 Hz) with 50 dB per octave slopes, and presented to two groups of normal-hearing listeners. The first group of listeners rated the similarity of the filtered vowels to their unfiltered counterparts. The second group identified the vowels that they heard through phonetic transcription while noting any distortions. Results indicated that the optimal one-octave filter bands identified for these vowels allowed for the correct perception and identification of each vowel. These optimal octaves agreed with Guberina's findings (1964, 1972) and McKenney (1971), and McKenney and Asp (1972).

Guberina conducted similar experiments with consonants as well as vowels for all languages taught at the Institute of Phonetics at the University of Zagreb in Zagreb, Yugoslavia. British English is taught as a foreign language. The optimal octaves identified for British English phonemes by Guberina are illustrated on Table IV. Further, Roberge (1979) in his work with phonetic correction has compiled the optimal octaves for French, English, German, Italian,

Table IV. Optimal Octaves for British English Consonants and Vowels

Optimal Octaves	Consonants Optimal	Vowels English (British)
100-200 Hz	-	
150-300 Hz	-	
200-400 Hz	p b	u uə
300-600 Hz	w	ʊ ɒ
400-800 Hz	l n	ə ɹ
600-1200 Hz	k g f v m	ɑ ʌ ɔ
800-1600 Hz	t d h l	ɑ ə ɔ i
1200-2400 Hz	ʃ ʒ r d ʒ	æ
1600-3200 Hz	ʃ tʃ	ɪ e ɛ
2400-4800 Hz	ʃ	eɪ aɪ
3200-6400 Hz	θ	i
4800-9600 Hz	z	
6400-12800 Hz	s	

Source: Institute of Phonetics, University of Zagreb, Zagreb, Yugoslavia, from Guerbina (1964, 1972).

and Spanish phonemes from other studies. Similarly, optimal octaves were determined in a Verbo-tonal Workshop held at The University of Tennessee, Knoxville (Asp and Guberina, 1982) and were used in the remediation of articulation deficits of normal-hearing children (Spletzer and Asp, 1981). The slopes for the one-octave filter bands were 20-26 dB per octave. Silverstein (1984) determined the optimal octaves with a 11-14 dB per octave slopes. He has used these optimal octaves for correcting misarticulations and foreign dialects.

III. Use of Optimal Octaves in Speech Detection Testing

Based on his work involving optimal octaves in phoneme perception, Guberina (1964) developed speech stimuli for Verbo-tonal Audiometry. The speech stimuli consisted of consonants and vowels with similar or adjacent optimal one-octave filter bands (Asp, 1975). Further, they are homogeneous with respect to the pitch model (Asp et al., 1977). Table V illustrates the nonsense syllables along with their identified optimal octave and comparable pure tones. These nonsense syllables are called "logotomes" and are repeated twice to simulate the normal rhythm and linguistic redundancy of speech (Asp, 1975).

These logotomes in both the filtered and unfiltered forms are used for speech detection testing in Verbo-tonal Audiometry for obtaining air and bone conduction thresholds. Guberina and Asp (1981)

Table V. Logotomes Used in Verbo-Tonal Audiometry with Their Optimal Octaves and Comparable Pure Tone

Logotomes	Optimal Octave (Hz)		Comparable Pure Tone (Hz)
	Series A	Series B	
/brubru/	50-100	37.7-75	75
/mumu/	75-150	100-200	125
/bubu/	150-300	200-400	250
/vovo/	300-600	400-800	500
/la α la/	600-1,200	800-1,600	1,000
/keke/	1,200-1,400	1,600-3,200	2,000
/ʃiʃi/	2,400-4,800	3,200-6,400	3,000
/sisi/	4,800-9,600	6,400-12,800	6,000

Source: Guberina (1964, 1972).

explain that the patient's detection thresholds for the filtered and unfiltered logotomes are compared for his thresholds for pure tones. If the patient's threshold for the filtered logotome is substantially better than his threshold for the logotome's comparable pure tone, the patient may have the potential to discriminate speech better through the particular frequency region than would otherwise indicate by the pure tone threshold. However, if the patient's thresholds for the logotome is poorer than for the comparable pure tone, the patient may have difficulty using this particular frequency region for speech discrimination. Guberina has designed "transfer tests" using filter/unfiltered logotomes to assess patients' potential for perceptually transferring speech sounds as well as for assessing progress in auditory training. Asp (1975) and Guberina and Asp (1981) explain these tests in detail.

IV. Summary

This chapter has presented research in the areas of: 1) the pitch of speech sounds (Boring, 1942; Harbold, 1958; McGlone and Manning, 1975; Peterson, 1971; Peterson and Asp, 1972; Asp, 1975; Asp, Berry, and Bessell, 1977; Bessell, 1979; and Asp and Johnson, 1985); 2) optimal octaves in phoneme perception (Chiba and Kajiyama, 1958; Guberina, 1964, 1972; McKenney, 1971; McKenney and Asp, 1972; Miner and Danhauer, 1977; Roberge, 1979; Spletzer and Asp, 1981; and Silverstein, 1984); and 3) use of optimal octaves in speech detection testing (Guberina, 1964). From this presentation of the literature,

several observations can be made: 1) the research supports the notion that phonemes, syllables, and words can be ranked according to pitch; 2) there are optimal octaves for perception of phonemes with practical applications for both assessment and remediation.

CHAPTER III

METHODS

I. Listeners

Listeners were 12 normal-hearing and -speaking adults (10 female, two male) ranging in age from 23 to 40 years of age ($\bar{X}$ =29.75 years), from various geographic regions of the United States and Canada, with a wide range of clinical experience in speech-language therapy. Appendix A includes the sex, age, amount of clinical experience, occupation, geographic place of origin, and approximate years of residence in East Tennessee for each listener.

All listeners passed a hearing screening test for air conduction thresholds bilaterally at 20 dB HL for the octave frequencies 250-8,000 Hz with a standard clinical audiometer (Madsen, Model No. 0B-802), and had normal articulation.

All listeners had successfully completed a phonetics course and a minimum of one year (four academic quarters) of clinical practicum experience in speech-language therapy.

II. Speech Stimuli and Ordering Procedure

Speech stimuli consisted of five logotomes (Guberina, 1964, 1972): 1) /mumu/; 2) /vovo/; 3) /lɑlɑ/; 4) /keke/; and 5) /sisi/. These five logotomes were chosen because they represent the frequency range used in Verbo-tonal Audiometry (see Table VI), and the phonemes

are homogeneous with respect to the pitch model (Asp, Berry, and Bessell, 1977) (see Table II, page 12). For example, /mumu/ is from the low pitch (tonality) category, /vovo/ is from the low-middle category, /lɔlɔ/ is from the middle category, /keke/ is from the middle-high category, and /sisi/ is from the high category.

Two procedures were used for the order of presentation of:

1) rating the five unfiltered homogeneous nonsense syllables according to pitch within a paired-comparison paradigm; and 2) phonetic transcription and quality ratings of the same five nonsense syllables each filtered through six one-octave filter bands.

The order of presentation for the rating of the five unfiltered homogeneous nonsense syllables according to pitch within a paired-comparison paradigm was determined by a method suggested by Guilford (1952). Using the equation $(n(n-1))/2$, the experimenter calculated the number of possible pairs ($N=10$) to be judged for this given set of five nonsense syllables. According to Guilford (1952), it is not necessary for listeners to judge both orders of a pair (e.g. /mumu - sisi/ and /sisi - mumu/). Either order is permissible. However, the listeners in this study judged both orders of the ten pairs of nonsense syllables for a total of 20 pairs of nonsense syllables. Although this was not necessary, this procedure served as an intra-subject reliability check for the rating of both orders of a pair of nonsense syllables on pitch. Appendix B displays the presentation order of the rating of the five homogeneous nonsense

syllables on pitch within a paired-comparison paradigm. In addition, the asterisked pairs of nonsense syllables represent those orders of the pairs that will be used for all calculations of intra- and inter-subject reliability as well as statistical calculations for determining the ranking of the five nonsense syllables on pitch. These ten asterisked items were randomly selected. In addition, a random presentation order of ten practice pairs of stimuli not included in the experiment was used to train the listeners.

The order of presentation for the phonetic transcription and quality ratings of the five homogeneous nonsense syllables each filtered through six one-octave filter bands was determined by randomizing three lists of ten nonsense syllables with each of the five nonsense syllables appearing twice per list. Appendix B displays the presentation order for the phonetic transcription and quality ratings of the five homogeneous nonsense syllables each filtered through six one-octave filter bands. In addition, a random presentation order of ten nonsense syllables not used as stimuli in the actual experiment was used to train the listeners on these tasks.

III. Recording Procedure for the Speech Stimuli

The speech stimuli were spoken by an experienced adult male speaker of General American English who was originally from the Northeastern section of the United States, but who has spent the last 30 years of his life in Knoxville, Tennessee.

Each nonsense syllable was repeated twice (e.g. /mumu/) to simulate the normal rhythm and linguistic redundancy of speech (Asp, 1975). The speech stimuli were recorded in a sound-treated booth (IDE, Model No. 130) in the exact order determined by the procedures described above. The speaker sat at a table inside the sound-treated booth and spoke approximately six inches away from the microphone (Unisphere, Model No. 5655D) for an optimal signal-to-noise ratio. The output of the microphone was routed outside the sound-treated booth to the microphone input of a reel-to-reel tape recorder (Revox, Model No. B77). Inside the sound-treated booth was a timing device to signal the speaker in seven and one-half second intervals. Outside of the sound-treated suite sat the experimenter to insure that the input of the speech stimuli to the tape recorder did not exceed -5 dB SPL on the VU meter of the tape recorder (see Figure 2).

The speaker read the speech stimuli from a list during recording. He was instructed to articulate the nonsense syllables with "natural effort" allowing for: 1) a 0.5 second silent interval within members of a pair of unfiltered nonsense syllables and a six second silent interval between pairs of nonsense syllables (e.g. /sisi/ - (0.5 second silent interval) - /mumu/ (six second silent interval) /la la/ - (0.5 second silent interval) - /keke/) for the rating of pitch; and 2) a 6.5 second silent interval between nonsense syllables filtered through six one-octave filter bands (e.g. /mumu/ (6.5 second silent interval) /sisi/ for phonetic transcription and quality ratings.

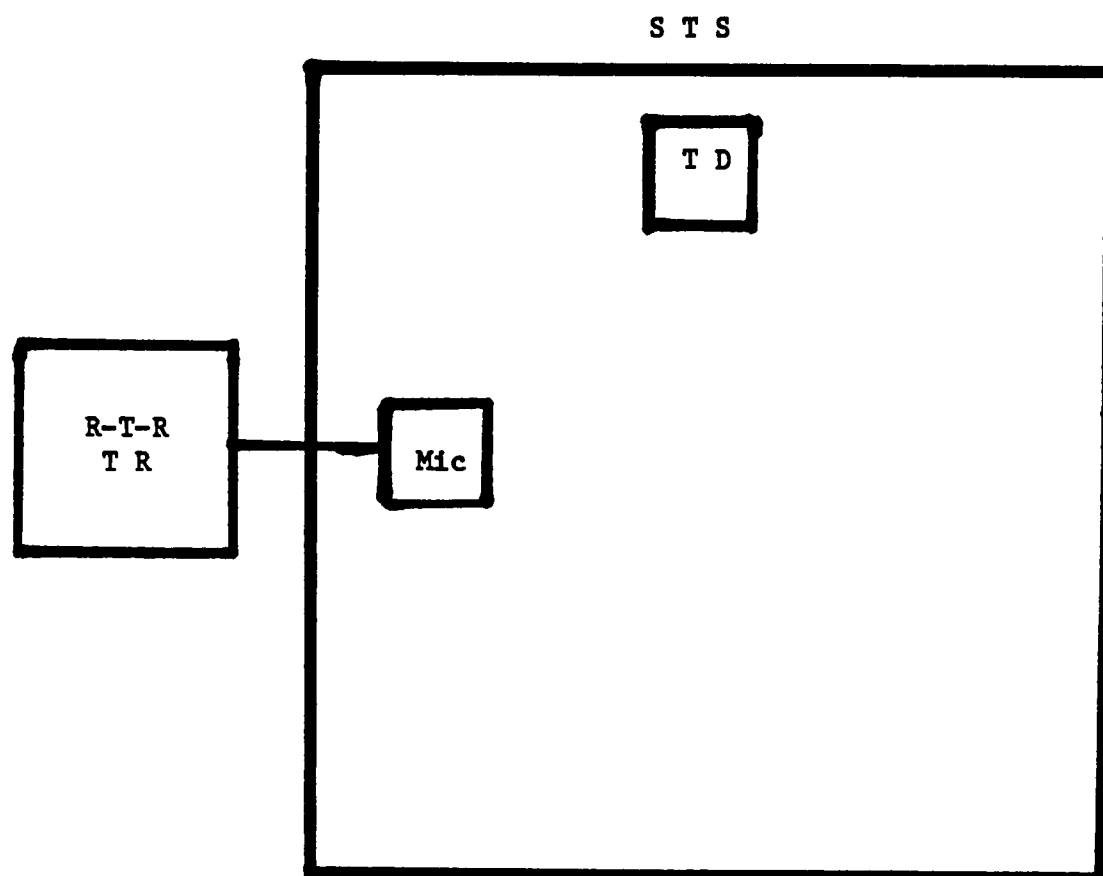


Figure 2. Block Diagram for Recording of Speech Stimuli.

STS = Sound-Treated Suite; R-T-R/TR = Reel-to-Reel Tape Recorder;
Mic = Microphone; TD = Timing Device.

These recorded speech stimuli were played individually to three speech-language pathologists with Certificates of Clinical Competence who unanimously judged the recordings to be: 1) intelligible; 2) relatively noise-free; and 3) undistorted reproductions of the original speech stimuli.

IV. Instrumentation

The SUVAG (Systems Universal Verbo-tonal Audition Guberina) Lingua was used as the filtering instrument in this study. Guberina and Asp (1981) described this multi-filter unit as having seven independent channels: 1) a flat frequency response (0.5 to 20,000 Hz, electrically); 2) a low-pass filter of 320 Hz and lower frequencies; 3) a high-pass filter of 3,200 Hz and higher frequencies; and 4) four band-pass filters, each of which is capable of 1/3-, 1/2-, and 1-octave settings with a wide range of center frequencies (8 to 8,000 Hz). This unit is a high quality system that can be used for research and/or therapy.

The reel-to-reel tape recording of the speech stimuli was played on the same tape recorder that was used for recording the speech stimuli. The output of the reel-to-reel tape recorder was routed to the input of the SUVAG Lingua for filtering. The output of the SUVAG Lingua was routed to the input of two sets of Koss 6-A circumaural headsets. One set was used for the experimenter's monitoring of the speech stimuli while the other set of headsets were

for the listeners' participation in the experimental procedure (see Appendix C for frequency responses of the headsets).

The reel-to-reel tape recorder and the SUVAG Lingua were situated outside of the sound-treated suite (Suttles, Model No. SE 228) to minimize listener distraction during their participation in the experimental procedure. The experimenter communicated with the listeners from outside the sound-treated suite during the experiment via a simple microphone, amplifier, and loudspeaker set-up (see Figure 3).

V. One-Octave Bands and Presentation Levels

One-Octave Bands

Six one-octave filter bands with the following center frequencies were used in this study: 250, 500, 1,000, 2,000, 4,000, and 8,000 Hz. Appendix D displays the cut-off frequencies for each one-octave filter band. An electrical analysis of the SUVAG Lingua (Letowski, 1984) revealed that the one-octave band-pass filters used in this study had between a 20-26 dB per octave slope which was within manufacturer's specifications. Figure 4 illustrates four possible frequency responses of a band-pass filter of the SUVAG Lingua set on a center frequency of 1,000 Hz with: 1) 1L octave band setting with an 11-14 dB/octave slope; 2) a one-octave band setting with a 20-26 dB/octave slope; 3) a 1/2-octave band setting with a 36 dB/octave slope; and 4) 1/3-octave band with a 40 dB/octave slope.

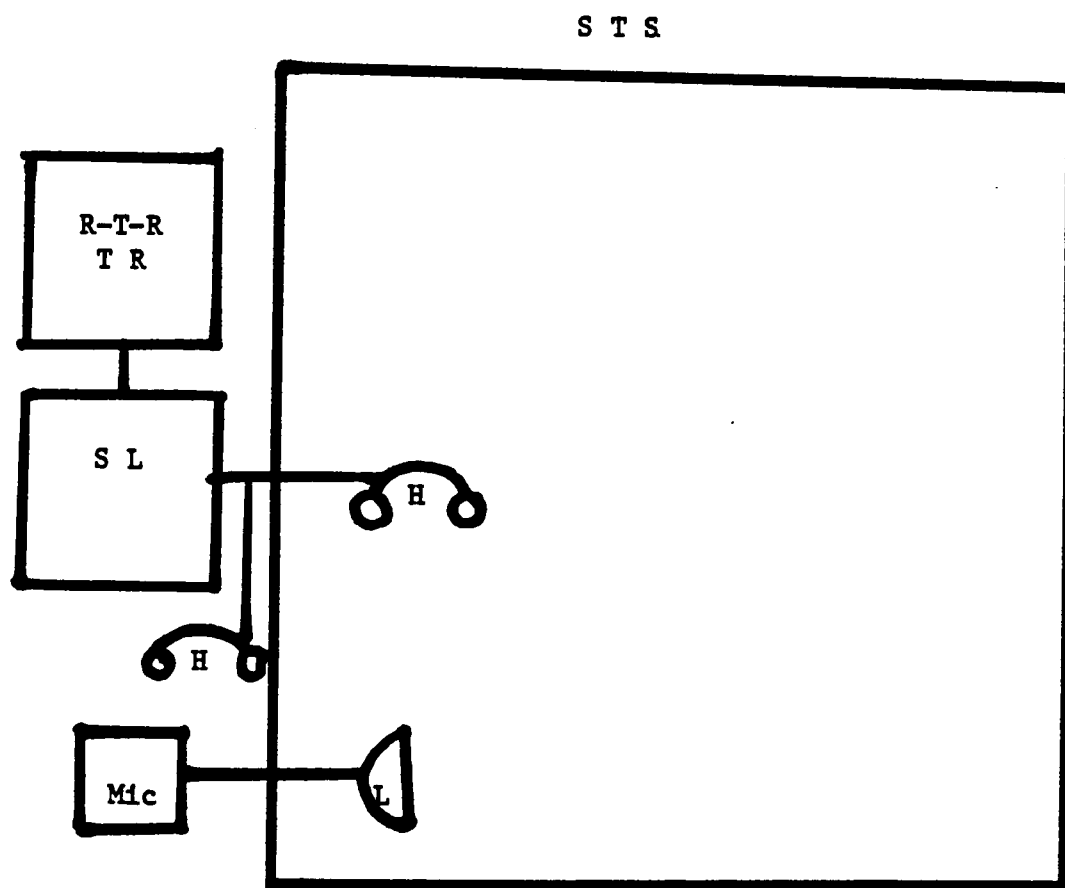


Figure 3. Block Diagram of Instruments for Experimental Sessions.

STS = Sound-Treated Suite; R-T-R/TR = Reel-to-Reel Tape Recorder;
SL = SUVAG Lingua; H = Headsets; Mic = Microphone; L = Loudspeaker.

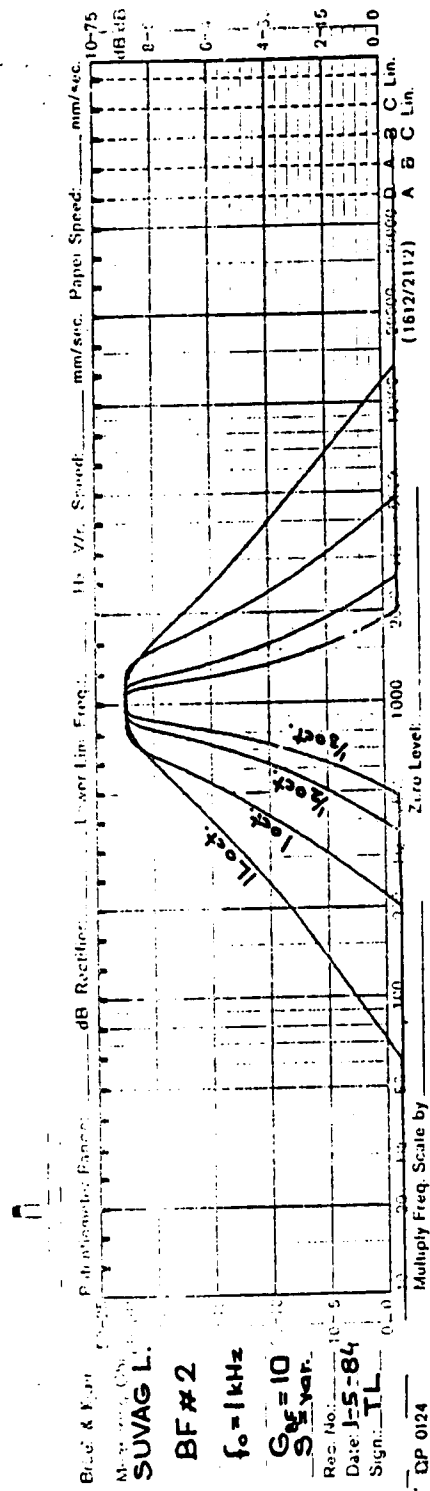


Figure 4. Frequency Responses of a Band-Pass Filter on a SUVAG Lingua at 1K Hz with a 1L Octave, One-Octave, 1/2-Octave, and 1/3-Octave Setting.

Source: Letowski (1984).

The six one-octave filter bands used in this study were selected because: 1) one-octave filter bands are the most frequently used band-pass filters in Verbo-tonal Audiometry; and 2) the center frequencies are closely related to the optimal octaves identified for the five logotomes used as speech stimuli in this study (see Table V).

Presentation Levels

For listener rating of the relative pitch of the five unfiltered homogeneous nonsense syllables, the SUVAG Lingua was set on a flat frequency response (0 to 20,000 Hz, electrically) and adjusted for a most-comfortable listening level. This level was approximately 70 dB SPL at the listeners' ears via Koss K-6A headsets.

For the phonetic transcription and quality ratings of the same five homogeneous nonsense syllables filtered through six one-octave filter bands, two presentation modes were used: 1) the "attenuator not adjusted"; and 2) the attenuator adjusted for "relative equal loudness."

For the "attenuator not adjusted" presentation mode, the attenuator of the SUVAG Lingua was kept constant with only the center frequencies of the one-octave filter bands being changed during presentation of the speech stimuli. This presentation mode resulted in some nonsense syllables being louder than others. For example, the nonsense syllable /sisi/ sounded much louder than /mumu/ when filtered through

the 8,000 Hz one-octave filter band. This phenomenon was created by the filter acting on the different spectra of each nonsense syllable.

For "relative equal loudness" presentation mode, a panel of two experienced listeners adjusted the attenuator of the band-pass filter of the SUVAG Lingua with all other controls remaining the same as the other presentation mode in order to determine the relative equal loudness of the five homogeneous nonsense syllables each filtered through six one-octave filter bands. These relative equal loudness levels are plotted for each nonsense syllable filtered through six one-octave filter bands in reference to the presentation mode of "no attenuator adjusted" (see Appendix E).

Both presentation modes resulted in listening levels that were all audible enough to complete the task, but some listening levels were louder than others. Future research is needed to determine the effect of differential loudness levels on listeners' phonetic transcriptions and quality ratings of homogeneous nonsense syllables each filtered through six one-octave filter bands.

VI. Listening Sessions

All listening sessions took place in the same sound-treated with the instruments arranged as described above (see section IV. Instrumentation and Figure 3). Listeners were tested one at a time

in the sound-treated suite. Each listener was given a brief introduction to the experimental tasks consisting of three parts: 1) phonetic transcription of nonsense syllables filtered through one-octave filter bands; 2) rating the pitch of unfiltered nonsense syllables; and 3) assessing the quality of filtered nonsense syllables. While the experimenter prepared the equipment, each listener was given a phonetic transcription "warm-up" exercise to complete. Upon completion of the exercise, any questions were answered by the experimenter.

During the listening session, each listener was given three different rating booklets, in succession, corresponding to each of the three parts of the experiment. Each listener was given the same instructions and training procedure for each of the three parts of the experiment. Appendix F displays the experimental protocol which contains the instructions and training procedure for each part of the experiment. However, for Part III -- quality ratings of the five homogeneous nonsense syllables filtered through six one-octave filter bands -- additional instruction and training was needed than what appears on the Part III response booklet protocol in Appendix F. In particular, listeners needed additional training in using the five-point quality rating scale. It was necessary to strictly define for the listeners and play examples of what ratings on the scale would be. In general, a one on the rating scale was defined as low quality or "unintelligible"; a three was average quality of "intelligible" (i.e. what was heard in normal conversation); and five was

high quality or "enhanced perception" (i.e. better perception than 'real life' -- a "clinical model").

Speech stimuli were presented in the prior determined orders described earlier (see Section II. Speech Stimuli and Ordering Procedure, page 25, and Appendix B) with the experimenter making the necessary adjustments on the SUVAG Lingua during the experiment. In the following sequence, each listener: 1) identified through phonetic transcription the five homogeneous nonsense syllables filtered through six one-octave filter bands for both the "attenuator not adjusted" and "relative equal loudness" presentation modes; 2) rated the pitch of the five unfiltered homogeneous nonsense syllables within a paired-comparison paradigm; and 3) assessed the quality of the five homogeneous nonsense syllables filtered through the same six one-octave filter bands for the "relative equal loudness" presentation mode.

At the completion of the experiment, each listener completed a post-experiment questionnaire eliciting demographic information and opinions concerning: 1) difficulty of the experimental tasks; 2) clarity of the signal; 3) presence of any distractions; 4) degree of motivation to complete the tasks; 5) confidence in responses; and 6) how the experimental tasks could be improved to maximize concentration and performance (see Appendix E).

With a precision sound level meter (Bruel and Kjaer, Model No. 2209), the amount of ambient noise inside the sound-treated suite was measured as 16.5 dB (A) and 48.4 dB (C). In addition, octave

band measures were also taken and the results appear in Appendix G. All octave band measures were well within the limits allowable for audiometric testing booths (ANSI, 3.1-1977).

VII. Spectrographic Analysis

Spectrographical analyses were completed on all speech stimuli used in the study. For the analyses, the output of the reel-to-reel tape recorder (Revox, Model No. B77) was routed to the input of the SUVAG Lingua for amplification and/or filtering. The output of the SUVAG Lingua was routed to the tape input of a digital sona-graph (Kay Elemetrics, Model No. 7800). The output of the digital sona-graph was connected to the input of the sona-graph printer (Kay Elemetrics, Model No. 7900) (See Figure 5). Wide-band spectrographic analyses to 8,000 Hz (300 Hz filter) were completed on: 1) the five homogeneous nonsense syllables each filtered through six one-octave filter bands; and 2) the ten pairs of unfiltered homogeneous nonsense syllables. In addition, narrow-band analysis expanded to 2,000 Hz (45 Hz filter) were also completed on the ten pairs of unfiltered nonsense syllables. All spectrograms appear in Appendix H.

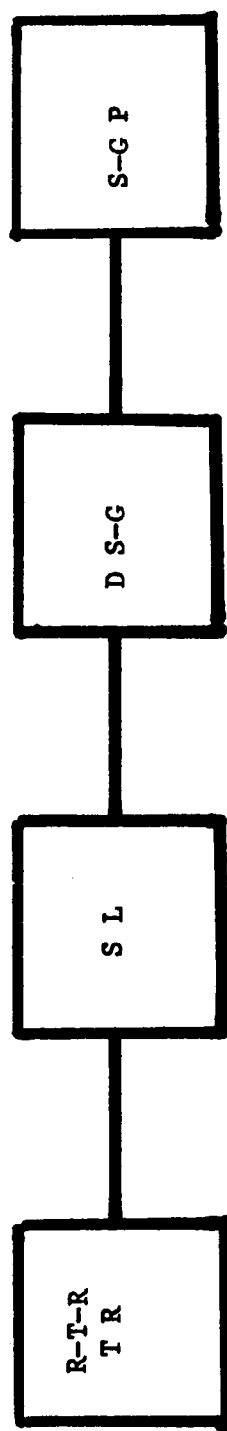


Figure 5. Block Diagram of Instruments for Spectrographic Analysis of Speech Stimuli.

R-T-R/TR = Reel-to-Reel Tape Recorder; SL = SUVAG Lingua; DS-G = Digital Sona-Graph;
S-GP = Sona-Graph Printer.

CHAPTER IV

RESULTS AND DISCUSSION

I. Post-Experiment Questionnaire and Reliability

Post-Experiment Questionnaire

Following completion of the experiment, each listener completed a post-experiment questionnaire eliciting listeners' demographic information and opinions concerning: 1) difficulty of the experimental tasks; 2) clarity of the speech stimuli; 3) presence of any distractions; 4) their degree of motivation to complete the tasks; and 5) their confidence in their responses.

Table VI illustrates the results of the post-experiment questionnaire. Results indicated that 100% (12/12) of the listeners: 1) felt that the experimental tasks were not difficult; 2) were motivated to complete the tasks; 3) reported that the speech stimuli was clear and that there were no distractions during the experiment; whereas 4) only 66% (8/12) were confident of their responses.

Reliability

Pitch (Tonality) Judgments

Intra-subject reliability. To calculate intra-subject reliability, each subject's responses were transferred to a matrix. Although listeners rated the pitch (tonality) of both orders of a

Table VI. Results of the Post-Experiment Questionnaire for 12 Listeners

Question	Percent	
	Yes	No
1. Difficulty of Task	100%	0%
2. Clarity of Signal	100%	0%
3. Presence of any Distractions	0%	100%
4. Motivation to Complete the Task	100%	0%
5. Confidence in Responses	66%	34%

pair (e.g. /sisi-mumu/ and /mumu-sisi/ to assess position effects on pitch judgments, only one order of each pair was used to determine intra-subject reliability (see explanation in Chapter III, Methods; Section II, Speech Stimuli and Ordering Procedure).

Figure 6 illustrates the matrix used in calculating intra-subject reliability for subject 1. The conversion of subjects' responses to matrices provide a visual display of pitch judgments making inconsistencies more evident. Listeners' rating inconsistencies determined intra-subject reliability.

In Figure 6, the order of the nonsense syllables are in the same order for both the ordinate and abscissa. A "1" in the coordinate of row X and column Y indicated that the nonsense syllable in row X was rated higher in pitch than the one in column Y. Further, an "0" must be placed in the coordinate corresponding to row Y and column X. For example, subject 1 rated "/si/" higher than "/mu/." Thus, "/si/" row and "/mu/" column coordinate had a "1". Then, a "0" was placed at the coordinate of row "/mu/" and column "/si/." Peterson (1971) explained the use of the matrices in determining intra-subject reliability. Appendix I contains the matrices for the other nine listeners. Subject 4 was unable to respond in the forced-choice response format of the paired-comparison paradigm. She was excluded from this task, but she did complete the other sections of the study.

In the determination of intra-subject reliability, it was necessary to calculate the number of circular triads for each listener

Matrix for Responses to Subject 1

	mu	vo	la	ke	si		ai	ai ²
mu	-	0	1	0	0		1	1
vo	1	-	0	0	0		1	1
la	1	1	-	0	0		1	1
ke	1	1	1	-	0		3	9
si	1	1	1	1	-		4	16
							Σ 10	Σ 28

Figure 6. Matrix of Responses for Subject 1.

Note: A "1" means that the nonsense syllable in row X was rated higher in pitch than the nonsense syllable in column Y. A "0" means that the nonsense syllable in row X was perceived as lower in pitch than the nonsense syllable in column Y.

using an equation formulated by Kendall (1955). Figure 7 illustrates an example of both a consistent ranking and an inconsistent ranking or circular triad. For example, if /sisi/ was rated higher in pitch than /lɔlɔ/, (i.e., /sisi → lɔlɔ/), /lɔlɔ/ higher than /mumu/ (i.e., /lɔlɔ → mumu/), and /sisi/ higher than /mumu/ (i.e., /sisi → mumu/), then the three nonsense syllables could be ranked from high to low pitch (i.e., /sisi → lɔlɔ → mumu/). This ranking was consistent and no circular triad was present. However, if /sisi → lɔlɔ/ and /lɔlɔ → mumu/, but /mumu → sisi/, the judgments were inconsistent and a circular triad was present. The larger the number of circular triads, the more inconsistent are the pitch judgments of a subject. Figure 6 shows a circular triad of a subject in this study. Subject 1 rated /lɔlɔ/ higher than /vovo/ (i.e., /lɔlɔ → vovo/) and /vovo/ higher than /mumu/ (i.e., /vovo → mumu/), but /mumu/ higher than /lɔlɔ/ (i.e., /mumu → lɔlɔ/). This was a circular triad.

Kendall (1962) defined a measure of consistency by the number of circular triads, and reported it as the coefficient of consistency (range 0 to 1.0). The larger the value, the more consistent the ratings. In this study, the maximum number of circular triads per subject was five. Table VII lists the eleven subjects from the most consistent to the least consistent listener, based on the coefficient of consistency. For example, nine subjects had coefficients of consistency of 1.0 (no circular triads); one subject (No. 1) had .80 (one circular triad); and one subject (No. 11) had .60 (two circular

CONSISTENT RATINGS

sisi → lala

lala → mumu

sisi → mumu

CIRCULAR TRIAD

sisi → lala

lala → mumu

mumu → sisi



Figure 7. Examples of Consistent and Inconsistent (Circular Triad) Ratings.

Table VII. Subjects' Number, Number of Circular Triads, and Coefficient of Consistency

Subject Number	Number of Circular Triads	Coefficients of Consistency
2	0	1.0
3	0	1.0
5	0	1.0
6	0	1.0
7	0	1.0
8	0	1.0
9	0	1.0
10	0	1.0
12	0	1.0
1	1	.80
11*	2	.60

* Subject 11 was eliminated from this portion of the experiment because coefficient of consistency was below .80. Also, subject 4 did not complete the part of the experiment.

triads). Other studies with the paired-comparison paradigm (Peterson, 1971; Peterson and Asp, 1972; Asp, Berry, and Bessell, 1978; and Bessell, 1979) used only the most consistent listeners. Therefore, only the subjects with coefficients of consistency of .80 or greater were included. Thus, the responses of subject 11 were not included in this part of the experiment. The consistency of the remaining 10 subjects was considered satisfactory for the experiment. Peterson (1971) explained the calculation of the coefficient of consistency.

Inter-subject reliability. To determine inter-subject reliability (agreements), the responses of each subject were transferred to a matrix (10 subjects x 10 pairs of nonsense syllables = 100 total responses recorded on the composite matrix). Figure 8 illustrates the composite matrix. The "ai" for each row is the number of times each nonsense syllable was rated higher in pitch out of a possible 40 times. The " Σ^2 " is the sum of the square of each cell number in each row. The "ai," " Σ^2 ," and " Σai^2 " values were used to calculate inter-subject reliability.

The coefficient of agreement among subjects was calculated with an equation and the number of circular triads (Kendall, 1962). It could range from 0 to 1.0. A larger value indicated more consistency among the subjects. The coefficient of agreement was $\pm .92$. This suggested excellent agreement among the listeners and was satisfactory for the purposes of this experiment. Peterson (1971) explained the calculation of the coefficient.

Composite Matrix of Responses

	mu	vo	la	ke	si	'	ai	y ²	ai ²
mu	-	1	1	0	0		2	2	4
vo	9	-	0	0	0		9	81	81
la	9	10	-	0	0		19	181	361
ke	10	10	10	-	0		30	300	900
si	10	10	10	10	-		40	400	1,600
							$\sum ai$	$\sum y^2$	$\sum ai^2$
							100	964	2,946

Figure 8. Composite Matrix of all Subjects' Ratings of the Five Homogeneous Nonsense Syllables on Pitch.

Reliability of both orders of a pair. The listeners judged both orders of pairs of the nonsense syllables (e.g., /sisi-mumu/ and /mumu-sisi/). Ninety-four percent (94%) of the ten listeners had the same judgment for both orders. This suggested that the position of nonsense syllables within the pair had little effect on the judgments. The listeners were considered reliable in their judgments.

Optimal Octaves Judgments

Intra-subject reliability. One subject was randomly selected to repeat 20% of the items on the identification and quality assessment tasks. An intra-subject reliability check revealed an 85% reliability for the identification task and 80% for the quality assessment task. These figures were calculated by dividing the number of agreements between items completed for the first vs. second time by the total number of items checked for agreement, and multiplied by 100. Intra-subject reliability of 85% and 80% were satisfactory for the purposes of this experiment.

Inter-subject reliability. The responses of two subjects on 20% of the items of the identification and quality assessment tasks were compared. Inter-subject reliability was 85% for the identification task and 50% for the quality assessment task. These figures were calculated by dividing the number of agreements between the subjects by the total number of items, and multiplied by 100. The inter-subject reliability of 85% for the identification task was

satisfactory, but the 50% reliability for the quality assessment was too low. Further use of the rating scale used in the quality assessment task is needed to determine possible reasons for low inter-subject reliability.

II. Results and Discussion

Pitch (Tonality)

Experimental Question 1: Can the five unfiltered homogeneous nonsense syllables be ranked on pitch?

Results: In order to rank order of the five nonsense syllables from highest to lowest pitch, the nonsense syllables were listed according to the number of times each was rated higher than the other nonsense syllables. The "ai" values in the composite matrix were used for the rank order (see Figure 8). For example (see Table VIII), /sisi/ was ranked highest (40 out of 40), whereas /mumu/ was ranked lowest (2 out of 40). In Table VIII, the "ai" values were similar to the predicted values from the tonality model (Asp et al., 1977).

Experimental Question 2: Is there a significant difference among the ranks of the five unfiltered homogeneous nonsense syllables on pitch?

Results: David (1963) has a test of equality unique to the method of paired-comparison which is the non-parametric analogue of equality of treatment means in an analysis of variance. Refer to Peterson (1971) for an explanation of the calculation of this statistic.

Table VIII. Rankings of Pitch for the Five Unfiltered Homogeneous Nonsense Syllables and Times Rated Higher

Nonsense Syllable	Times Rated Higher	Rank Order	Predicted from Tonality Model	
			ai Values	Category
/sisi/	40	1	40	High
/keke/	30	2	30	Mid. High
/lala/	19	3	20	Mid.
/vovo/	9	4	10	Low-Mid.
/mumu/	2	5	0	Low

The significance level was $P < .05$ using a Chi-square significance table with $n-1$ degrees of freedom (Bruning and Kintz, 1977). The critical value was 9.5 with $P < .05$ and four degrees of freedom. The value obtained using David's (1963) procedure was $\chi^2 = 97.76$ which exceeded the set critical value. Thus, a significant difference existed among the ranks of the nonsense syllables on pitch.

Then, individual tests of significance were conducted between the ratings of adjacently ranked nonsense syllables. From the composite matrix of subjects' responses (see Figure 8), proportions of how many times one nonsense syllable was rated higher than its adjacently ranked nonsense syllable were obtained. For example (see Figure 8), /vovo/ was rated higher in pitch than /mumu/ by nine out of ten listeners while only one listener rated /mumu/ higher than /vovo/. From this information, a proportion was obtained. To determine the significance of the proportion, a non-parametric test was applied to the data (Bruning and Kintz, 1977; pp. 222-224). Results indicated that a significant difference of $z = 3.58$ (1.96; $P < .05$) between the rank of /mumu/ and /vovo/.

This same procedure was completed for the ranks of all adjacent nonsense syllables. Results indicated that significant differences exist between all adjacently ranked nonsense syllables. Thus, /sisi/ was ranked significantly higher in pitch than /keke/; /keke/ higher than /lala/; /lala/ higher than /vovo/; and /vovo/ higher than /mumu/.

Discussion

The finding that the ranks of the five nonsense syllables were significantly different according to pitch support the tonality model (Asp et al., 1977). The tonality model was developed on the research that identified the optimal octaves and pitch ranking of both consonants and vowels. The tonality model assigns consonants and vowels into low, low-middle, middle, middle-high, and high pitch categories. When words or syllables are "homogeneous," the phonemes are selected from the same or adjacent pitch categories. In this study, the nonsense syllables were homogeneous. Each nonsense syllable represented one of the five categories of the tonality model: /mumu/ represents the low; /vovo/ the low-middle; /lala/ the middle; /keke/ the middle-high; and /sisi/ the high category. The significant difference between the pitch ranks of each adjacent nonsense syllable suggests a significant difference between the five categories of the tonality model.

In using the tonality model, researchers homogeneously group phonemes in selecting frequency specific words. For example, a consonant and a vowel can be combined from the low pitch category to form the word "moo" (i.e., /mu/) or form the high pitch category for the word "she" (i.e., /ʃi/).

Asp, Berry, and Bessell (1978) and Bessell (1979) had normal-hearing listeners judge the relative pitch of 30 words in which ten words each were homogeneous with respect to the low, middle, and high categories of the tonality model. Statistical analysis of the

data indicated that the words could be ranked on pitch and grouped into significantly different low, middle, and high pitch categories. The results of the present study agrees with Asp, Berry, and Bessell (1978) and Bessell (1979) in that studies demonstrated statistically significant differences between categories of the tonality model.

Peterson (1971) and Peterson and Asp (1972) showed that normal-hearing listeners' judgments of the relative pitch of 23 pre-vocalic consonants paired with the neutral vowel / α /: 1) could be ranked according to pitch; and 2) significantly differentiated into low, low-middle, middle, middle-high, and high pitch categories agreeing with Asp, Berry, and Bessell (1978) and Bessell (1979). However, the results of the present study and Peterson (1971) and Peterson and Asp (1972) were the only demonstrating significant differences between all five categories of the tonality model.

The rank order of the vowels of the five nonsense syllables on pitch was in agreement with the work of earlier investigators as cited by Boring (1942), Harbold (1958), and McGlone and Manning (1975) who investigated the pitch of voiced and whispered vowels. Harbold (1958) and McGlone and Manning (1975) found that the pitch of phonemes are more dependent upon the resonance of the pharyngeal/oral cavity rather than fundamental frequency (F_0).

This is supported by the results of the present study in that the rank order of pitch of the nonsense syllables containing the /i/, / α /, and /u/ vowels have a direct positive relationship to their

second formants reported by Peterson and Barney (1952). However, no relationship exists between the rank order of pitch of the vowels and their fundamental frequencies (Peterson and Barney, 1952).

In summary, the results of this study were striking in terms of the significant differences revealed between ranking of pitch categories of the tonality model. Practical applications of the results might include formulating similar frequency specific nonsense syllables for use in an unfiltered nonsense syllable test of speech discrimination (Asp and Johnson, 1985).

Optimal Octave Judgments

Experimental Question 1: Are there "optimal" one-octave filter bands for the perception of each of the five homogeneous nonsense syllables resulting from normal-hearing listeners' identification through phonetic transcription of the five homogeneous nonsense syllables each filtered through six one-octave filter bands?

Results: Subjects' phonetic transcriptions of the five filtered homogeneous nonsense syllables were tallied from Part I of the experiment. Although listeners heard the nonsense syllables repeated twice (e.g. /mumu/), they transcribed a single consonant and a vowel (e.g. /mu/). Consonants and vowels were worth one point in the tally. Percentages were calculated for incidences of when each nonsense syllable was identified as the particular nonsense syllable when filtered through the six one-octave filter bands. These percentages

were calculated for both the "relative equal loudness" and "no attenuator adjusted presentation modes" (see Table IX and Figures 9-13).

The identification task did not reveal specific optimal one-octave filter bands for the nonsense syllables. However, the nonsense syllables were perceived as some other consonant vowel combination or were too unintelligible for transcription when filtered through the 250 one-octave filter band.

Optimal center frequencies were estimated using a linear averaging technique of the center frequencies with adjacent bands that had identification scores of 96% or greater for each nonsense syllable. Only the "no attenuator adjusted" presentation mode was used for the identification scores because it seems to differentiate among some of the nonsense syllables. For example, /vovo/ was identified as /vovo/ 100% of the time when filtered through the 500, 1,000, and 2,000 Hz one-octave filter bands for an estimated mean optimal center frequency of 1,167 Hz. This is an estimate, because the task was too easy for these three conditions. If it was more difficult, one of these octave bands would have had a higher score and averaging would not have been necessary. The same procedure was completed for all nonsense syllables resulting in the following estimated mean optimal center frequencies: 1) /lɔlɔ/, 750 Hz; 2) /keke/, 3,750 Hz; and 3) /sisi/, 3,750 Hz. These mean center frequencies are linear and do not take into consideration the logarithmic relationships that listeners may use in pitch perception.

Table IX. Percent of Identification for Five Nonsense Syllables
Filtered Through Six One-Octave Filter Bands

	Filter Bands					
	250	500	1,000	2,000	4,000	8,000
/mumu/						
I	67%	96%	88%	96%	88%	96%
II	59%	92%	88%	92%	92%	92%
/vovo/						
I	71%	100%	100%	100%	96%	71%
II	50%	92%	100%	100%	100%	62%
/la-la/						
I	42%	100%	100%	92%	92%	88%
II	67%	100%	100%	100%	85%	100%
/keke/						
I	79%	96%	100%	100%	100%	100%
II	50%	98%	100%	100%	100%	100%
/sisi/						
I	75%	96%	100%	100%	100%	100%
II	55%	83%	100%	100%	100%	100%

Key: I = "No Attenuator Adjusted" presentation mode.

II = "Relative Equal Loudness" presentation mode.

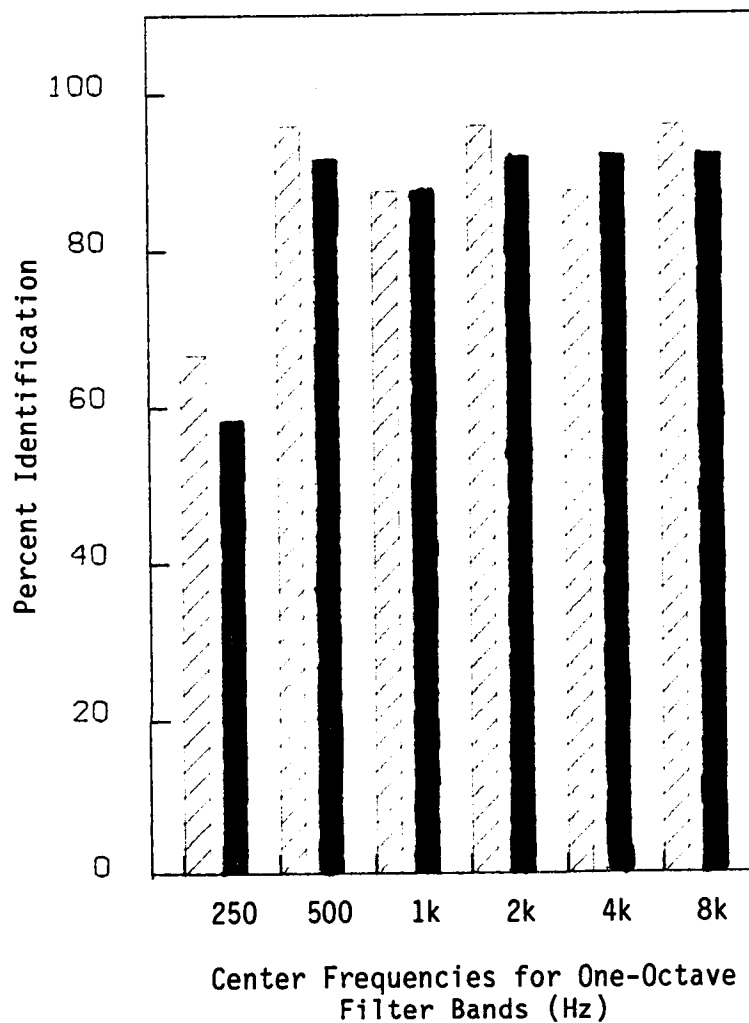


Figure 9. Bar Graph of Identification Scores for "No Attenuator Adjusted" (Cross-hatched) and "Relative Equal Loudness Modes" (Solid) for /mumu/ Filtered Through Six One-Octave Filter Bands.

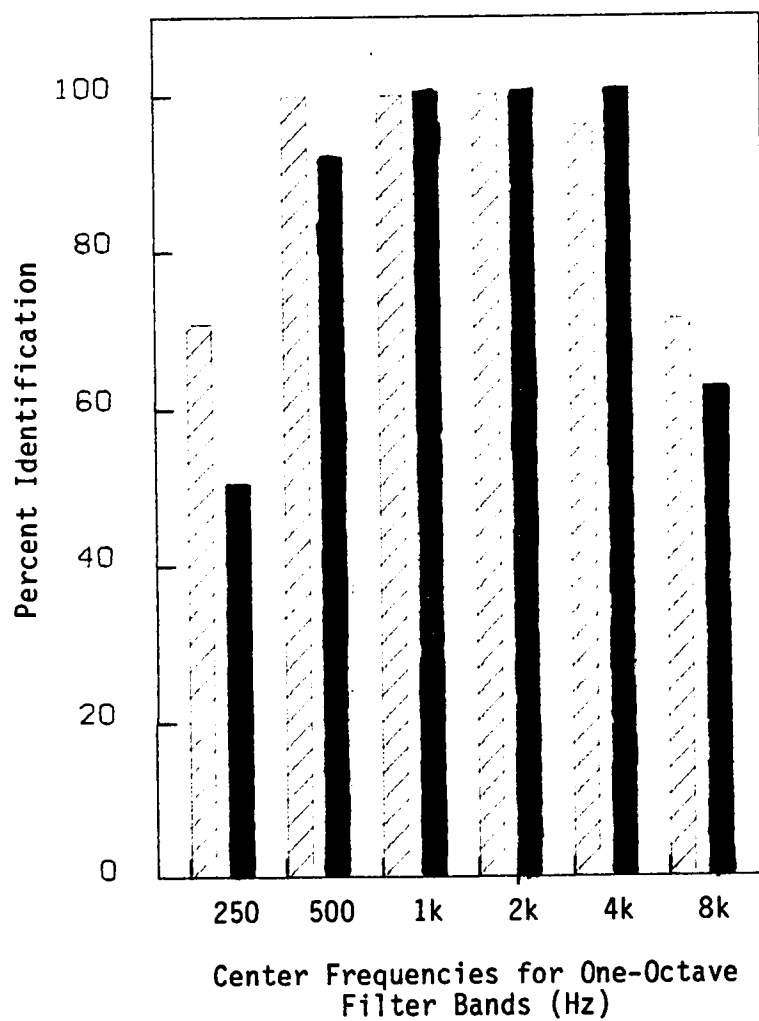


Figure 10. Bar Graph of Identification Scores for "No Attenuator Adjusted" (Cross-hatched) and "Relative Equal Loudness Modes" (Solid) for /vovo/ Filtered Through Six One-Octave Filter Bands.

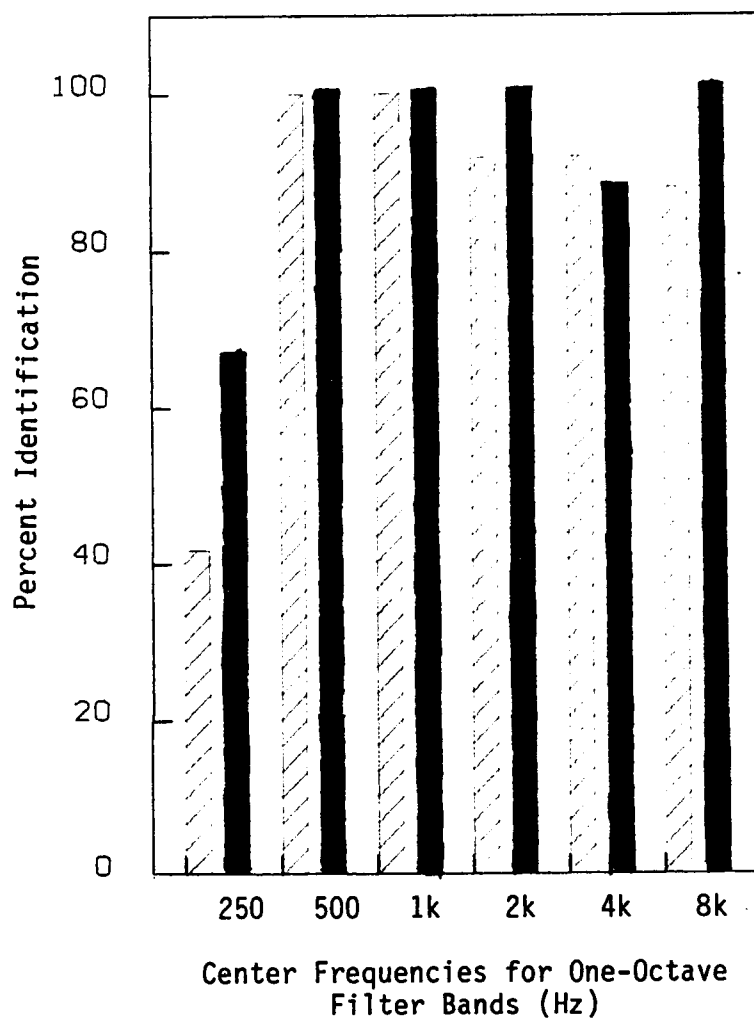


Figure 11. Bar Graph of Identification Scores for "No Attenuator Adjusted" (Cross-hatched) and "Relative Equal Loudness Modes" (Solid) for $/l\alpha l\alpha/$ Filtered Through Six One-Octave Filter Bands.

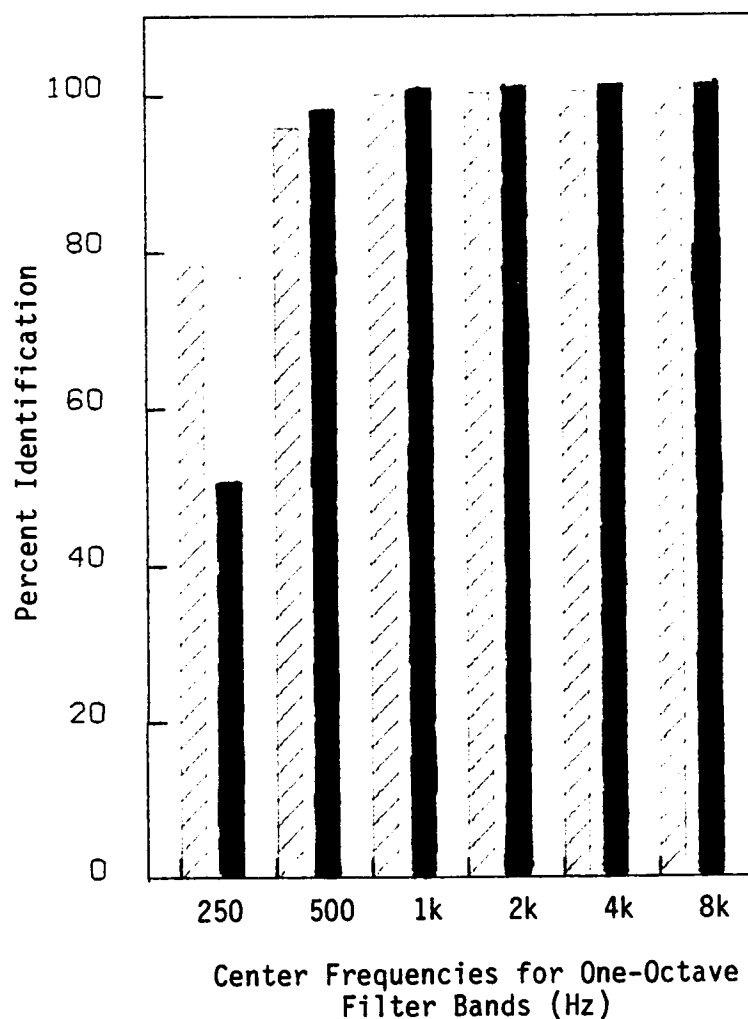


Figure 12. Bar Graph of Identification Scores for "No Attenuator Adjusted" (Cross-hatched) and "Relative Equal Loudness Modes" (Solid) for /keke/ Filtered Through Six One-Octave Filter Bands.

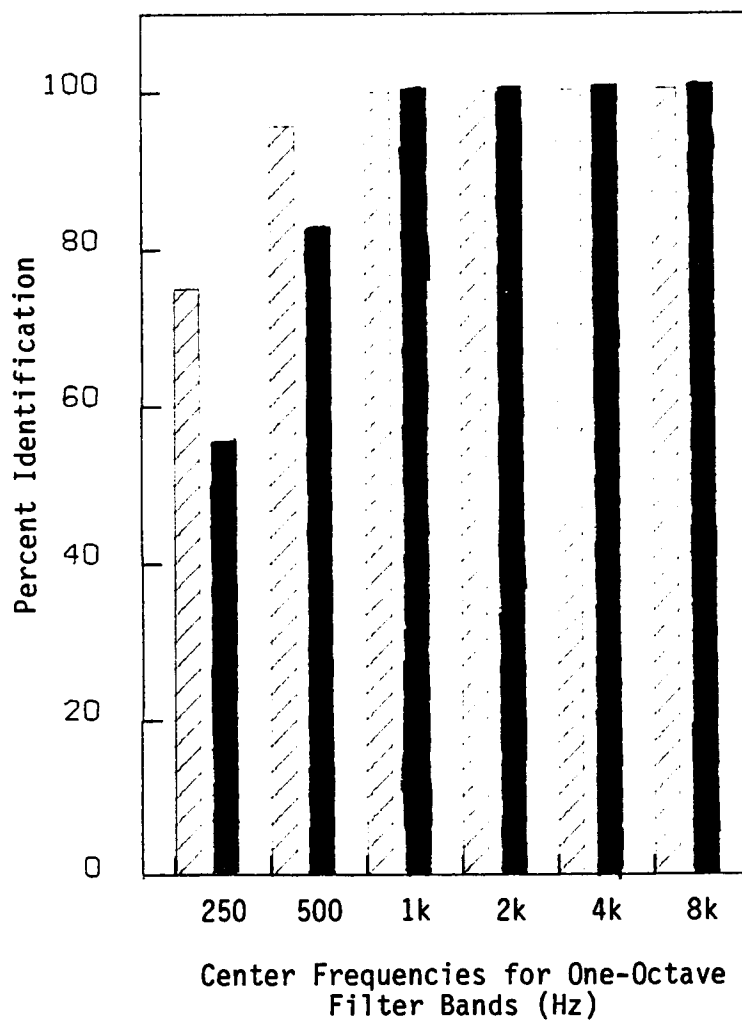


Figure 13. Bar Graph of Identification Scores for "No Attenuator Adjusted" (Cross-hatched) and "Relative Equal Loudness Modes" (Solid) for /sisi/ Filtered Through Six One-Octave Filter Bands.

A mean optimal center frequency was not calculated for /mumu/ because the bands of 500, 2,000, and 8,000 Hz were not adjacent. Other researchers identified two optimal octaves for the nasals. For example, Roberge (1979) reported two optimal octaves for /m/ of 600-1,200 plus 1,200 and 2,400 Hz. Thus, the present study used 500 Hz plus 2,000 or 8,000 Hz to represent the two bands of nasals.

Discussion: The identification task was too easy for the listeners as evidenced by 100% identification scores for some nonsense syllables in more than one octave band. One hundred percent identification scores did not differentiate among bands. Thus, specific optimal one-octave filter bands were not identified. Only the 250 band had any deleterious effect on the intelligibility of the nonsense syllables.

Possible contributing factors to the ease of the task included: 1) a closed-response type task created by the presence of only five nonsense syllables in the experiment; 2) the repetition of each nonsense syllable to simulate the normal rhythm and linguistic redundancy of speech (Asp, 1975). The nonsense syllables repeated twice may be as intelligible as words. Furthermore, the gradual slope of 20-26 dB per octave may have contributed to the identification of nonsense syllables outside of their optimal frequency regions. For example, /mumu/ was identified as /mumu/ through the 8,000 Hz one-octave filter band. The estimated mean optimal center frequencies are tabled and discussed later.

Experimental Question 2: Are there "optimal" one-octave filter bands for the perception of each of the five homogeneous nonsense syllables resulting from the normal-hearing listeners' quality ratings of the five homogeneous nonsense syllables each filtered through six one-octave filter bands?

Results: Subjects' quality ratings of the five filtered homogeneous nonsense syllables were judged on a scale of "1" low quality to "5" high quality on Part III of the experiment. From these judgments, the mean quality ratings and standard deviations were calculated (see Table X and Figures 14-18).

In Table X, the octave bands with the asterisked values had the highest scale values for each nonsense syllable. These highest values were significantly different from other bands based on an adjacent t-test procedure used by Miner and Danhauer (1977) (see Tables XI-XV). This procedure required identifying the optimal one-octave band with the highest mean quality rating and systematically comparing it to all other bands with a t-test with repeated measures (Bruning and Klintz, 1977). For /mumu/ and /keke/ one optimal one-octave filter band was identified for each one. However, for /vovo/, /la la/, and /sisi/, more than one band was identified. In such cases, estimated optimal center frequencies were calculated using the same linear averaging technique used for the identification task. The estimated mean optimal center frequencies for both tasks appear in Table XVI.

Table X. Means and Standard Deviations of the Quality Ratings of Five Homogeneous Nonsense Syllables Filtered Through Six One-Octave Bands

	Center Frequency of Octave Filter Bands (in Hz)					
	250	500	1,000	2,000	4,000	8,000
<i>/mumu/</i>						
$\bar{X}$ =	1.33	2.41	2.5	3.58*	2.0	1.83
SD =	.49	.9	.67	.9	0	.58
<i>/vovo/</i>						
$\bar{X}$ =	1.08	2.08	3.41*	3.75*	3.0	1.75
SD =	.29	.79	.79	.75	.74	.45
<i>/la la/</i>						
$\bar{X}$ =	1.0	3.25*	3.58*	3.91*	3.08	1.91
SD =	0	1.05	1.0	1.08	1.44	.51
<i>/keke/</i>						
$\bar{X}$ =	1.08	2.00	3.58	4.25*	3.33	3.08
SD =	.29	.43	.67	.96	.98	.79
<i>/sisi/</i>						
$\bar{X}$ =	1.16	1.75	2.5	3.91*	4.25*	3.91*
SD =	.39	.45	.52	.90	.75	.99

*Significantly higher quality ratings (by rows).

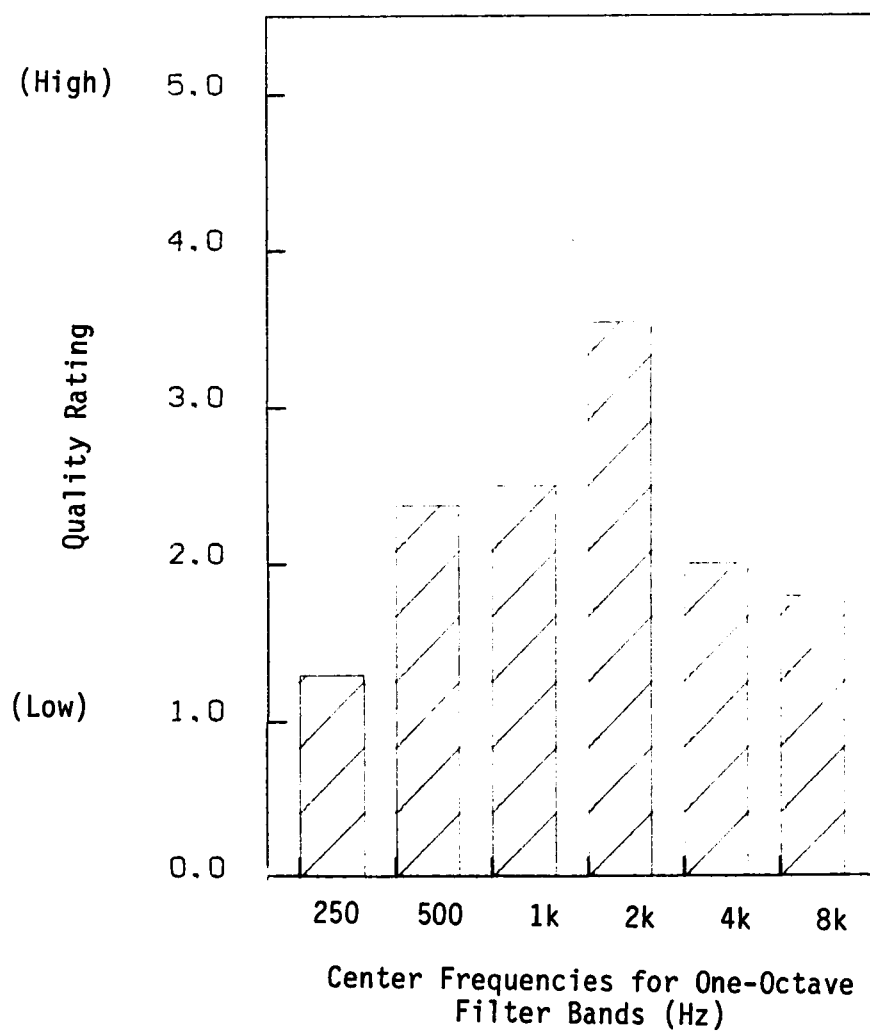


Figure 14. Bar Graph of Quality Ratings of /mumu/ Filtered Through Six One-Octave Filter Bands.

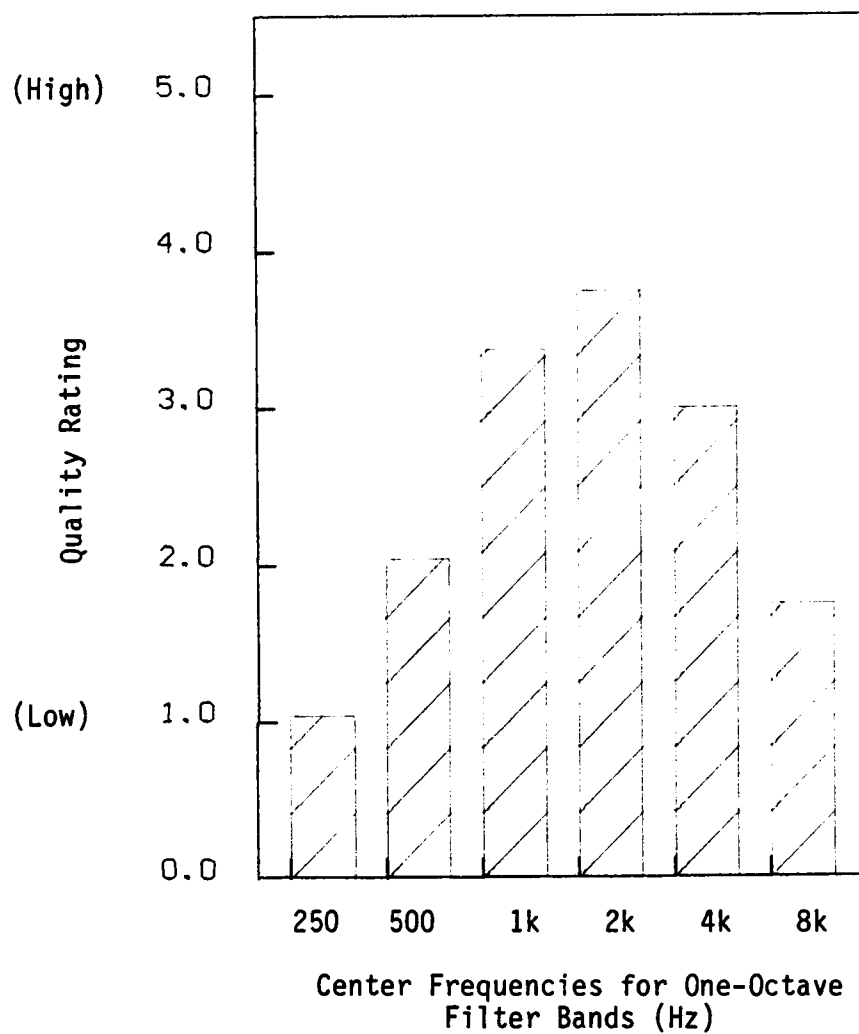


Figure 15. Bar Graph of Quality Ratings of /vovo/ Filtered Through Six One-Octave Filter Bands.

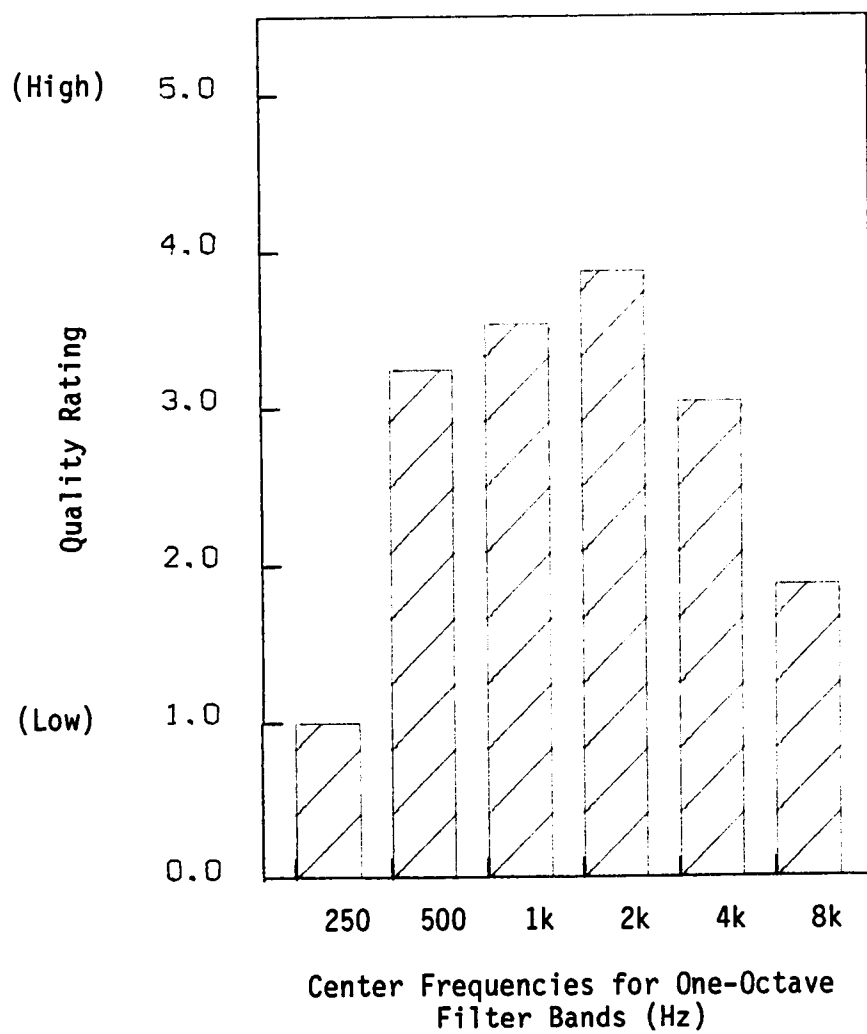


Figure 16. Bar Graph of Quality Ratings of /1010/ Filtered Through Six One-Octave Filter Bands.

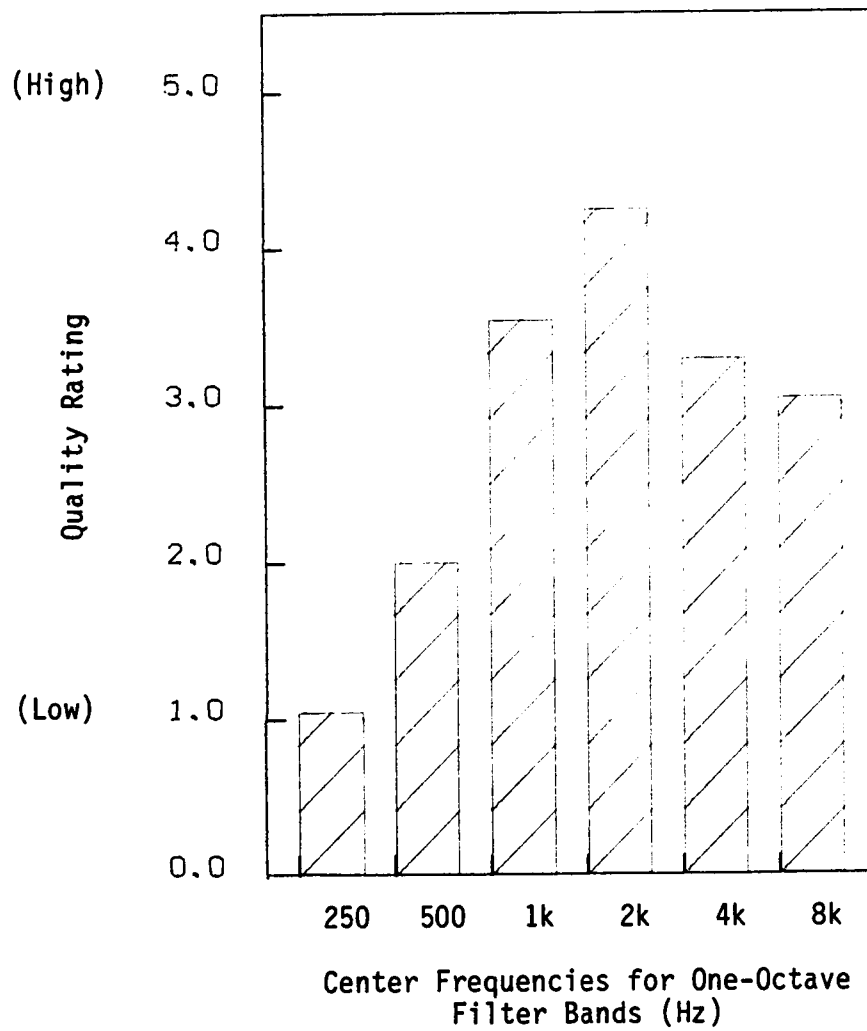


Figure 17. Bar Graph of Quality Ratings of /keke/ Filtered Through Six One-Octave Filter Bands.

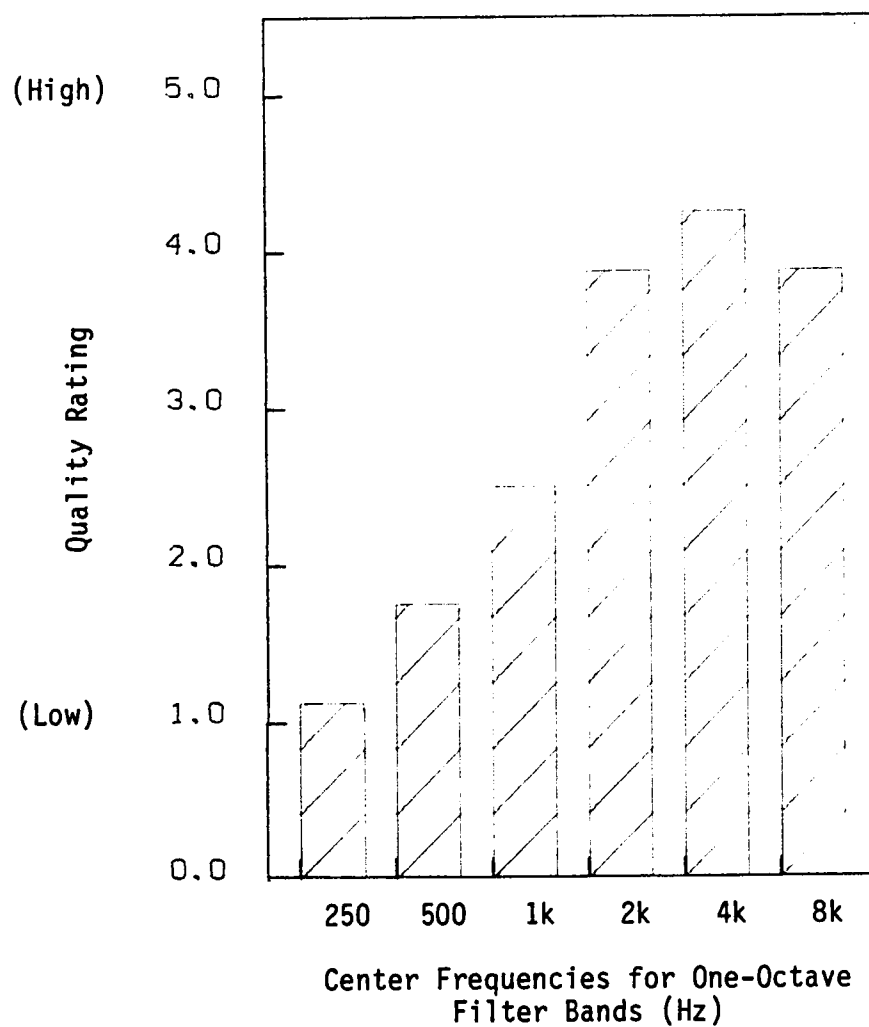


Figure 18. Bar Graph of Quality Ratings of /sisi/ Filtered Through Six One-Octave Filter Bands.

Table XI. Adjacent t-Test Procedure for Obtaining Optimal Frequency Regions for the Nonsense Syllable /mumu/*

Octave Band Center Frequencies	t-Score**	p	Significant?
250 vs. 2,000	5.79	< .0005	Yes
500 vs. 2,000	3.2	< .005	Yes
1,000 vs. 2,000	3.78	< .005	Yes
2,000 vs. 4,000	6.34	< .0005	Yes
2,000 vs. 8,000	7.0	< .0005	Yes

*Subjects' quality ratings of /mumu/ filtered through six one-octave filter bands revealed 2,000 Hz as the optimal one-octave band.

**With 11 degrees of freedom.

Table XII. Adjacent t-Test Procedure for Obtaining Optimal Frequency Regions for the Nonsense Syllable /vovo/*

Octave Band Center Frequencies	t-Score**	p	Significant?
250 vs. 2,000	12.5	< .0005	Yes
500 vs. 2,000	5.44	< .0005	Yes
1,000 vs. 2,000	1.32	> .05	No
2,000 vs. 4,000	3.0	< .01	Yes
2,000 vs. 8,000	9.52	< .0005	Yes

*Subjects' quality ratings of /vovo/ filtered through six one-octave filter bands revealed 1,000 and 2,000 Hz as the optimal one-octave bands.

**With 11 degrees of freedom.

Table XIII. Adjacent t-Test Procedure for Obtaining Optimal Frequency Regions for the Nonsense Syllable /la₁la/ *

Octave Band Center Frequencies	t-Score**	p	Significant?
250 vs. 2,000	7.4	< .0005	Yes
500 vs. 2,000	1.13	> .05	No
1,000 vs. 2,000	.84	> .10	No
2,000 vs. 4,000	2.15	< .05	Yes
2,000 vs. 8,000	8.0	< .0005	Yes

*Subjects' quality ratings of /la₁la/ filtered through six one-octave filter bands revealed 500, 1,000, and 2,000 Hz as the optimal one-octave bands.

**With 11 degrees of freedom.

Table XIV. Adjacent t-Test Procedure for Obtaining Optimal Frequency Regions for the Nonsense Syllable /keke/*

Octave Band Center Frequencies	t-Score**	p	Significant?
250 vs. 2,000	11.74	< .0005	Yes
500 vs. 2,000	10.22	< .0005	Yes
1,000 vs. 2,000	1.91	< .05	Yes
2,000 vs. 4,000	3.02	< .01	Yes
2,000 vs. 8,000	3.9	< .005	Yes

*Subjects' quality ratings of /keke/ filtered through six one-octave filter bands revealed 2,000 Hz as the optimal one-octave band.

**With 11 degrees of freedom.

Table XV. Adjacent t-Test Procedure for Obtaining Optimal Frequency Regions for the Nonsense Syllable /sisi/*

Octave Band Center Frequencies	t-Score**	p	Significant?
250 vs. 2,000	11.45	< .0005	Yes
500 vs. 2,000	6.54	< .0005	Yes
1,000 vs. 2,000	4.53	< .0005	Yes
2,000 vs. 4,000	1.32	> .05	No
4,000 vs. 8,000	1.03	> .05	No

*Subjects' quality ratings of /sisi/ filtered through six one-octave filter bands revealed 2,000, 4,000, and 8,000 Hz as the optimal one-octave bands.

**With 11 degrees of freedom.

Table XVI. Estimated Mean Optimal Center Frequencies for the Five Homogeneous Nonsense Syllables (in Hz)

Nonsense Syllable	Johnson (1986)		Guberina (1964/1972)			Asp & Guberina (1982)		
	Identification	Quality (Mean Center Frequency)	Consonant and Vowel	Linear Center Frequency	Cut-off Frequencies	Linear Center Frequency	Cut-off Frequencies	
/mumu/	500 plus (2000 or 8000)	2,000	/m/	900	600-1,200	800	165-1,131	
/vovo/	1,167	1,500	/u/	300	200-400	250	176-353	
/lɔlɔ/	750	1,167	/v/	900	600-1,200	640	452-905	
			/o/	600	400-800	500	353-707	
/keke/	3,750	2,000	/l/	600	400-800	500	353-707	
			/ɔ/	1,200	800-1,600	1,250	880-1,760	
/sisi/	3,750	4,000	/k/	900	600-1,200	1,250	880-1,760	
			/e/	2,400	1,600-3,200	2,000	1,411-2,820	
			/s/	9,600	6,400-12,800	8,000	5,650-11,310	
			/i/	4,800	3,200-6,400	5,000	3,530-7,070	

Discussion: Table XVI displays Guberina's (1964; 1972) and Asp and Guberina's (1982) optimal octaves for both consonants and vowels. The center frequencies were calculated by linearly averaging the cut-off frequencies for the phonemes. Note that the optimal octave bands reported by Guberina (1964; 1972) for the individual consonant and vowels were not the same as the optimal octaves he reported for the entire logotome for Verbotonal Audiometry (see Table VI, page 41). At this time, it is not clear how Guberina (1964; 1972) obtained the optimal octaves for the logotomes. Guberina and Asp (1981) stated that the logotomes' optimal octave bands for Verbotonal Audiometry were those bands in which the logotomes were most easily identified by normal-hearing listeners.

The estimated mean center frequencies for the individual phonemes were not averaged in obtaining a single estimated optimal center frequency for each logotome. Future research is needed to determine if an averaging technique can be used to predict the optimal octaves of syllables and words from those of individual phonemes. Nevertheless, the estimated mean optimal center frequencies obtained in this study were in fairly good agreement with Guberina's (1964; 1972) and Asp and Guberina's (1982) optimal octaves for the individual phonemes for /la|a/, /keke/, and /sisi/. However, the estimated mean optimal center frequencies obtained for /mumu/ and /vovo/ were higher in frequency than for their individual phonemes (Guberina, 1964; 1972) and Asp and Guberina (1982). Future research is needed to determine the nature of the discrepancy.

Further inspection of the data suggested that comparing estimated mean optimal center frequencies may not be the most valid method of analysis. Thus, Table XVII shows the cut-off frequencies identified in this study for the logotomes. These frequency regions were obtained by taking the lowest cut-off frequency of an optimal octave as well as the highest. In cases when only one optimal octave for a logotome was identified, the area between the cut-off frequencies represented the optimal frequency region. However, if two or more adjacent optimal octave bands were identified, the area between low cut-off frequency for the lowest adjacent band and the high cut-off frequency of the highest adjacent band represented the optimal frequency region of the logotome. Similarly, for Guberina's (1964; 1972) and Asp and Guberina's (1982) optimal octaves for the individual phonemes, the region between the low cut-off frequency of the vowel and the high cut-off frequency of the consonant represented the optimal frequency region for the logotomes predicted by this data.

Comparison of optimal frequency regions shows good agreement between the results of this study, and Guberina (1964; 1972) and Asp and Guberina (1982) for /vovo/, /la]a/, /keke/, and /sisi/. However, the optimal frequency region identified by /mumu/ in this study was higher in frequency than for optimal frequency region predicted by the optimals of /m/ and /u/ (Guberina, 1964, 1972; Asp and Guberina, 1982).

The relationship between the rank order of the five unfiltered nonsense syllables on pitch and their filtered optimal octaves was not

Table XVII. Optimal Frequency Regions for the Five Homogeneous Non-sense Syllables (in Hz)

NS	<u>Johnson (1986)</u>		Guberina (1964; 1972)	Asp & Guberina (1982)
	Identification	Quality		
/mumu/	500 Hz plus 2,000 or 8,000	1,410-2,820	200-1,200	353-1,131
/vovo/	707-2,620	707-2,820	400-1,200	353-905
/lala/	353-1,410	353-2,820	400-1,600	353-1,760
/keke/	707-11,310	1,410-2,820	600-3,200	880-2,820
/sisi/	707-11,310	1,410-11,310	3,200-12,800	3,530-11,310

clear. The two nonsense syllables /sisi/ and /keke/ with the highest optimal octaves or frequency regions were also highest ranked on pitch. However, the relationship was not as clear for /mumu/, /vovo/, and /lala/. Future research is needed.

CHAPTER V

SUMMARY AND CONCLUSIONS

I. Summary

No study has yet to simultaneously investigate the pitch of unfiltered speech stimuli and their optimal octaves for perception. Consequently, the purposes of this study were to have twelve normal-hearing listeners -- two males, ten females -- ages 23 to 40 years ($\bar{X} = 29.75$), 1) judge the relative pitch of five unfiltered nonsense syllables -- /mumu/, /vovo/, /lɔlɔ/, /keke/, and /sisi/ -- with the consonant and vowel homogeneously grouped according to a pitch model (Asp, 1975; Asp, Berry, and Bessell, 1977); and 2) identify through phonetic transcription and quality assessment of the same five homogeneous nonsense syllables each filtered through six one-octave filter bands with 20-26 dB per octave slopes and center frequencies of 250, 500, 1,000, 2,000, 4,000, and 8,000 Hz.

For the judgment of the five unfiltered homogeneous nonsense syllables on pitch, the listeners were presented pairs of nonsense syllables binaurally through headsets at a most comfortable listening level within a paired-comparison paradigm. Their task was to select which nonsense syllable was "higher" in pitch. Statistical analyses of the data revealed that: 1) the nonsense syllables could be ranked from high to low pitch -- /sisi/, /keke/, /lɔlɔ/, /vovo/, and /mumu/; and 2) there was a significant difference between adjacently ranked

nonsense syllables. Since each one of the nonsense syllables represented one of five pitch categories of the pitch model (Asp et al., 1977), these results indicate that a significant difference exists between each of the five pitch categories of the model.

For the identification of the five homogeneous nonsense syllables normal-hearing listeners phonetically transcribed each of the nonsense syllables filtered through six one-octave filter bands. The stimuli were presented binaurally through headsets at a loudness level sufficient to adequately complete the task.

From listeners' transcriptions of the filtered nonsense syllables, percentages were calculated for the incidence of each nonsense syllable being perceived as the nonsense syllable when filtered through each one-octave filter band. Results indicated high identification scores for all nonsense syllables. The nonsense syllables were most often perceived as some other nonsense syllable or too unintelligible for transcription when filtered through the 250 Hz one-octave filter band. Mean estimated optimal center frequencies were calculated by a linear averaging technique using the center frequencies of the adjacent one-octave filter bands with the highest identification score (i.e., 96-100%) for each nonsense syllable. These estimated optimal center frequencies were compared to Guberina's (1964; 1972) and Asp and Guberina's (1982) and were found to be in agreement for /lɒlɒ/, /keke/, and /sisi/, but not for /mumu/ or /vovo/. Clearly,

the identification task was too easy for listeners. This task needs to be replicated using steeper dB per octave slopes and a more difficult listening situation. Lower identification scores may make it easier to identify the "optimal" bands.

For listeners' quality assessment of the five homogeneous nonsense syllables each filtered through six one-octave filter bands, listeners used a five point rating scale ranging from "1" low quality to "5" high quality. Statistical analyses of the data revealed optimal center frequencies for all five nonsense syllables. The optimal center frequencies identified for /la_lla/, /keke/, and /sisi/ were in agreement with the optimal octaves identified by Guberina (1964; 1972). However, the optimal center frequencies identified for /mumu/ and /vovo/ were higher in frequency than for those identified by Guberina (1964; 1972) and Asp and Guberina (1982). Further inspection of the data revealed better agreement with Guberina (1964; 1972) and Asp and Guberina (1982) for both the results of the identification and quality assessment tasks when comparing optimal frequency regions rather than estimated optimal center frequencies.

In comparing the rank order of tonality of the five unfiltered nonsense syllables to the filtered optimal octaves, /sisi/ and /keke/ were ranked highest on pitch and had the highest optimal frequency regions. However, the relationship between tonality and optimal octaves of /mumu/, /vovo/, and /la_lla/ was not clear.

II. CONCLUSIONS

From the results of this study, it was concluded that:

1. The five homogeneous nonsense syllables were ranked from high to low pitch (tonality) -- /sisi/, /keké/, /lɔlɔ/, /vovo/, and /mumu/ -- based on normal-hearing listeners' judgments of relative pitch within a paired-comparison paradigm.

2. Significant differences were found among the all adjacently ranked nonsense syllables on pitch. Since /sisi/ was from the high pitch category, /keke/ from the middle-high, /lɔlɔ/ from the middle, /vovo/ from the low-middle, and /mumu/ from the low of the tonality model (Asp et al., 1977), significant differences existed between each of the pitch categories of the tonality model. These results agreed with the tonality model.

3. The identification task was too easy for the listeners. Estimated optimal center frequencies from the results of this task for /lɔlɔ/, /keke/, and /sisi/ were in agreement with Guberina's (1964; 1972) and Asp and Guberina's (1982) optimals, but not for /mumu/ or /vovo/.

4. The quality assessment task resulted in estimated optimal center frequencies for /lɔlɔ/, /keke/, and /sisi/ in agreement with Guberina's (1964; 1972) and Asp and Guberina (1982) optimals, but not for /mumu/ and /vovo/.

5. Good agreement between the results of this study and Guberina (1964; 1972) and Asp and Guberina (1982) was obtained when optimal

frequency regions were compared rather than estimated optimal center frequencies.

6. There was good agreement between the rank order of the pitch of the unfiltered /sisi/ and /keke/ and the order predicted from their optimal octaves or frequency regions, but this relationship was not clear for /mumu/, /vovo/, and /lɔlɔ/.

BIBLIOGRAPHY

BIBLIOGRAPHY

- Asp, C. W. (1972) The Verbo-tonal Method. Unpublished manuscript. Knoxville: The University of Tennessee.
- Asp, C. W. (1975) Measurement of aural speech perception and oral speech production of the hearing-impaired. In S. Singh (Ed.) Measurement Procedures in Hearing, Speech, and Language. Baltimore: University Park Press.
- Asp, C. W. (1985) The Verbo-tonal System. Graduate Course offered by the Department of Audiology and Speech Pathology at The University of Tennessee, Knoxville.
- Asp, C. W., Berry, J. S., and Bessell, C. S. (1977) A pitch (tonality) model for english phonemes. Unpublished Study: The University of Tennessee, Knoxville.
- Asp, C. W., Berry, J. S., and Bessell, C. S. (1978) The relative perceived pitch of 30 monosyllabic words: The rank order in comparison to a proposed pitch model. Journal of the Acoustical Society of America.
- Asp, C. W. and Guberina, P. (1982) The Optimal Octaves of English Phonemes. An unpublished study completed at the Third Verbo-tonal Summer Workshop held at the Western Pennsylvania School for the Deaf, Pittsburgh, PA.
- Asp, C. W. and Johnson, C. E. (1985) New applications and development of the tonality concept. Paper presented at the Third International Verbo-tonal Symposium and Workshop held in Zagreb, Yugoslavia.
- Bessell, C. S. (1979) Pitch Differences Among 30 English Words. Unpublished Master's Thesis, The University of Tennessee, Knoxville.
- Black, J. W. (1968) Perception of Altered Acoustic Stimulating the Deaf. Columbus: The Ohio State University Research Foundation.
- Boring, E. (1942) Sensation and Perception in the History of Experimental Psychology. New York: Appleton-Croft Co.
- Bruning, J. L. and Klintz, B. L. (1977) Computational Handbook of Statistics, 2nd Edition. Glenview, Ill.: Scott, Foresman and Company.

- Chiba, T. and Kajiyoma, M. (1958) The Vowel: It's Nature and Structure. Tokyo: Phonetic Society of Japan.
- David, H. A. (1963) The Method of Paired Comparisons. New York: Hafner Publishing Co.
- Donahue, A. M. (1983) Auditory sensitivity in the range of extra high frequencies for male and female populations. Unpublished Master's Thesis, The University of Tennessee, Knoxville.
- Guberina, P. (1964) Verbo-tonal Method: Its application to the rehabilitation of the deaf. Proceedings of the International Congress in Education of the Deaf. U.S. Government Printing Office, Washington, D.C.
- Guberina, P. (1972) Case studies in the use of restricted bands of frequencies in auditory rehabilitation of the deaf. Office of Vocational Rehabilitation, H.E.W. Contract OVR-YU60-2-63.
- Guberina, P. and Asp, C. W. (1981) The Verbo-tonal Method of Rehabilitating People with Communication Problems. Monograph 13 sponsored by the World Rehabilitation Fund, Inc., New York. National Institute for Handicapped Research. Grant No. 22-p-59037/2-03.
- Guilford, J. P. (1954) Psychometric Methods. 2nd Edition. New York: McGraw-Hill Press.
- Harbold, G. (1958) Pitch ratings of voiced and whispered vowels. Journal of the Acoustical Society of America. 30, 600(A).
- Kendall, M. G. (1955) Further contributions to the theory of paired comparisons. Biometrics, 11, 43-62.
- Kendall, M. G. (1962). Rank Correlation Methods. 3rd Edition. New York: Hafner Publishing Co.
- Koike, K. J. M., Brown, B., Hobbs, H., and Asp, C. W. (1983) Frequency selective speech discrimination test: Tennessee Tonality Test. Paper presented at the Annual Convention of the American Speech-Language-Hearing Association, Cincinnati, OH.
- Letowski, T. (1984) An electroacoustic analysis of SUVAG Auditory Training Units. An unpublished study, The University of Tennessee, Knoxville.
- McGlone, R. and Manning, W. (1975) Pitch of voiced and whispered vowels. Journal of the Acoustical Society of America. 58, S91.

- McKenney, E. E. (1971) Five English Vowels under Eleven Band-Pass Filtering Conditions as a Test of Optimal Octaves of Perception. Unpublished Master's Thesis, The University of Tennessee, Knoxville.
- McKenney, E. E. and Asp, C. W. (1972) Five English Vowels under eleven band-pass filtering conditions as a test of optimal octaves of perception. Journal of the Acoustical Society of America, 51, 122(A).
- Miner, R. and Danhauer, J. L. (1977) Relation between formant frequencies and optimal octaves in vowel perception. Journal of the American Audiological Society, 2, 163-168.
- Peterson, G. E. and Barney, H. L. (1952) Control methods used in a study of the vowels. Journal of the Acoustical Society of America, 24, 175-184.
- Peterson, M. M. (1971) A Study of the Perceived Pitch of 23 English Consonants. Unpublished Master's Thesis, The University of Tennessee, Knoxville.
- Peterson, M. M. and Asp, C. W. (1972) A study of the perceived pitch of 23 pre-vocalic consonants. Journal of the Acoustical Society of America, 52, 146(A).
- Poole, I. (1934) Genetic development of articulation of consonant sounds in speech. Elementary English Review, 11, 159-161.
- Roberge, C. (1979) Phonetic Correction and Language Teaching Theory and Practice of the Zagreb Language Teaching Method. Tokyo: Taishukan.
- Silverstein, B. (1984) The use of optimal frequencies in the correction of misarticulations and foreign dialects. Paper presented at the Fourth American Verbo-tonal Society Conference, Knoxville, TN.
- Spletzer, B. and Asp, C. W. (1981) The effectiveness of the Verbo-tonal System for normal-hearing children with speech and/or language problems. Final report of an investigation supported by a National Institute of Health Bio-Medical Science Support Grant from The University of Tennessee, Knoxville.
- Torgerson, W. (1958) Theory and Method of Scaling. New York: John Wiley and Sons.

APPENDIXES

APPENDIX A

SUBJECT DEMOGRAPHIC INFORMATION

SUBJECTS

<u>Subject Number</u>	<u>Sex</u>	<u>Age</u>	<u>Years Experience</u>	<u>Occupation</u>
1.	M	31	9	Spch. Path.
2.	F	40	17	Spch. Path.
3.	F	33	9	Spch. Path.
4.	F	37	10	Spch. Path.
5.	F	27	3	Spch. Path.
6.	F	28	2	Demo. Teacher
7.	F	28	6	Audiologist*
8.	M	32	2	Spch. Path.
9.	F	30	6	Audiologist*
10.	F	25	1	Spch. Path.**
11.	F	23	1	Spch. Path.**
12.	F	23	2	Spch. Path.**

* Extensive experience in Experimental Phonetics.

** Advanced Graduate Student at The University of Tennessee, Knoxville.

SUBJECTS

<u>Subject Number</u>	<u>Geographic Origin</u>	<u>Years in Knoxville</u>
1.	Arkansas	1
2.	Northern Calif.	1
3.	Eastern Kansas	4
4.	North Carolina	8
5.	East Tennessee	27
6.	North Carolina	3
7.	East Tennessee	28
8.	Ontario, Canada	2
9.	Pennsylvania	5
10.	East Tennessee	25
11.	New York	1
12.	Iowa	1

APPENDIX B

PRESENTATION ORDERS

PRESENTATION ORDER FOR THE IDENTIFICATION AND QUALITY ASSESSMENT
OF FIVE HOMOGENEOUS NONSENSE SYLLABLES FILTERED THROUGH
SIX-ONE-OCTAVE FILTER BANDS

Group I

	<u>Nonsense Syllable</u>	<u>Filter Band (Hz)</u>
1.	/mumu/	1,000
2.	/sisi/	4,000
3.	/vovo/	2,000
4.	/keke/	500
5.	/sisi/	250
6.	/la]a/	4,000
7.	/mumu/	250
8.	/vovo/	250
9.	/la]a/	500
10.	/keke/	2,000

Group II

1.	/keke/	4,000
2.	/mumu/	500
3.	/la]a/	1,000
4.	/vovo/	1,000
5.	/keke/	1,000
6.	/sisi/	8,000
7.	/la]a/	2,000
8.	/sisi/	1,000
9.	/vovo/	8,000
10.	/mumu/	2,000

Group III

1.	/la]a/	8,000
2.	/mumu/	4,000

	<u>Nonsense Syllable</u>	<u>Filter Band (Hz)</u>
3.	/vovo/	4,000
4.	/lɔlɔ/	250
5.	/keke/	8,000
6.	/sisi/	500
7.	/vovo/	500
8.	/sisi/	2,000
9.	/keke/	250
10.	/mumu/	8,000

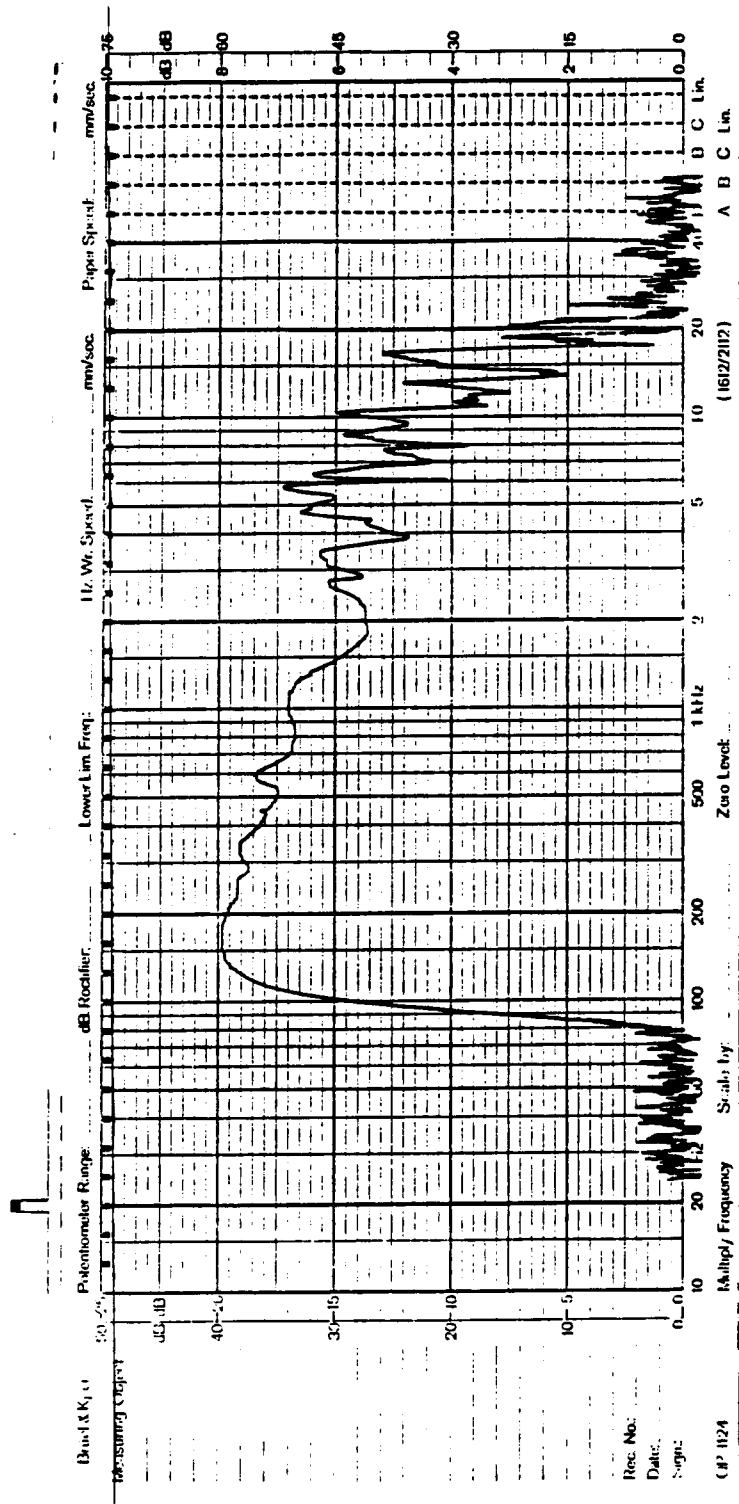
PRESENTATION ORDER FOR RATING THE PERCEIVED PITCH OF FIVE
UNFILTERED HOMOGENEOUS NONSENSE SYLLABLES

1. /sisi- vovo/*
2. /mumu- keke/*
3. /sisi- lɔlɔ/
4. /lɔlɔ- vovo/
5. /mumu- sisi/*
6. /vovo- lɔlɔ/*
7. /keke- mumu/
8. /sisi- mumu/
9. /mumu- lɔlɔ/*
10. /keke- sisi/*
11. /mumu- vovo/*
12. /vovo- keke/
13. /lɔlɔ- sisi/*
14. /keke- vovo/*
15. /lɔlɔ- keke/
16. /sisi- keke/
17. /vovo- mumu/
18. /keke- lɔlɔ/*
19. /lɔlɔ- mumu/
20. /vovo- sisi/

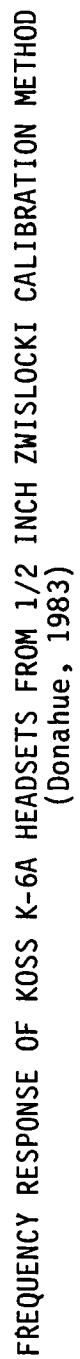
* Randomly selected pairs of stimuli to determine ranking of five unfiltered nonsense syllables on perceived pitch.

APPENDIX C

FREQUENCY RESPONSES OF THE HEADSETS



FREQUENCY RESPONSE OF KOSS K-6A HEADSETS FROM 1/4 INCH ZWISLOCKI CALIBRATION METHOD
 (Donahue, 1983)



APPENDIX D

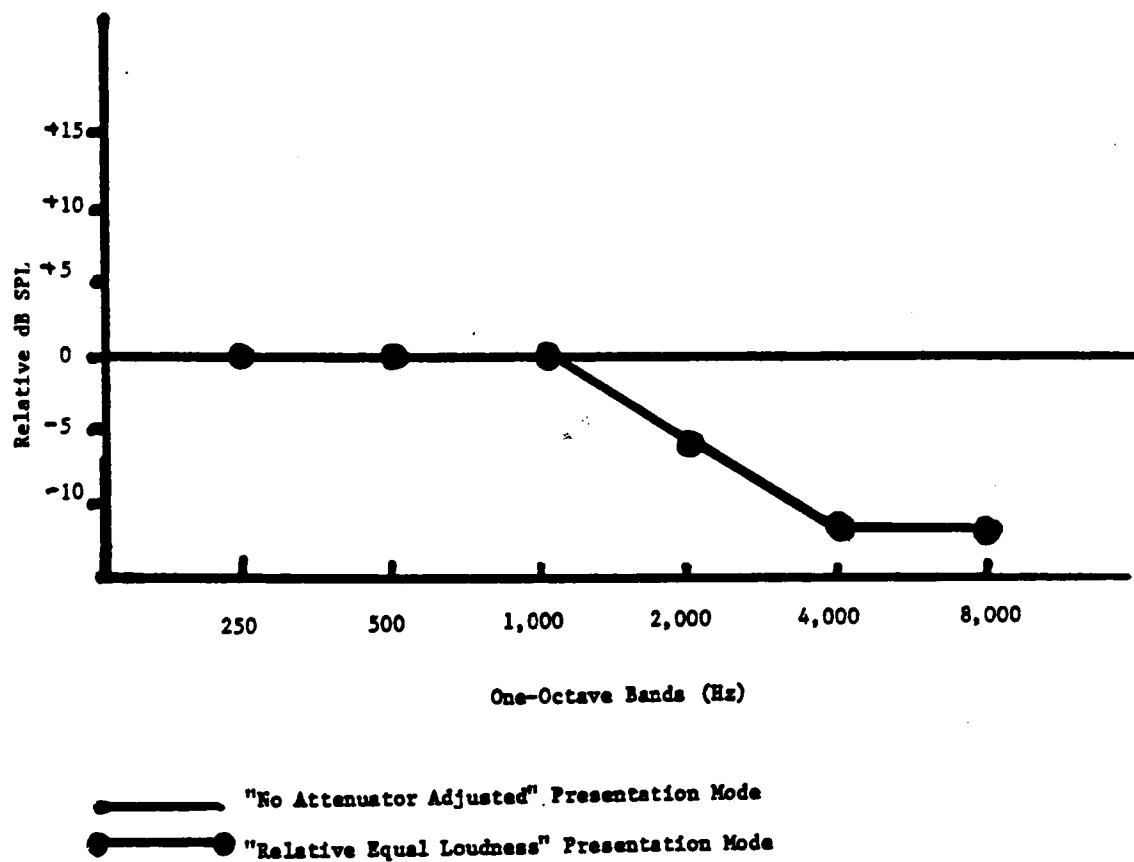
CENTER AND CUT-OFF FREQUENCIES OF THE ONE-OCTAVE FILTER BANDS

SUVAG LINGUA FOR INDIVIDUAL CORRECTION

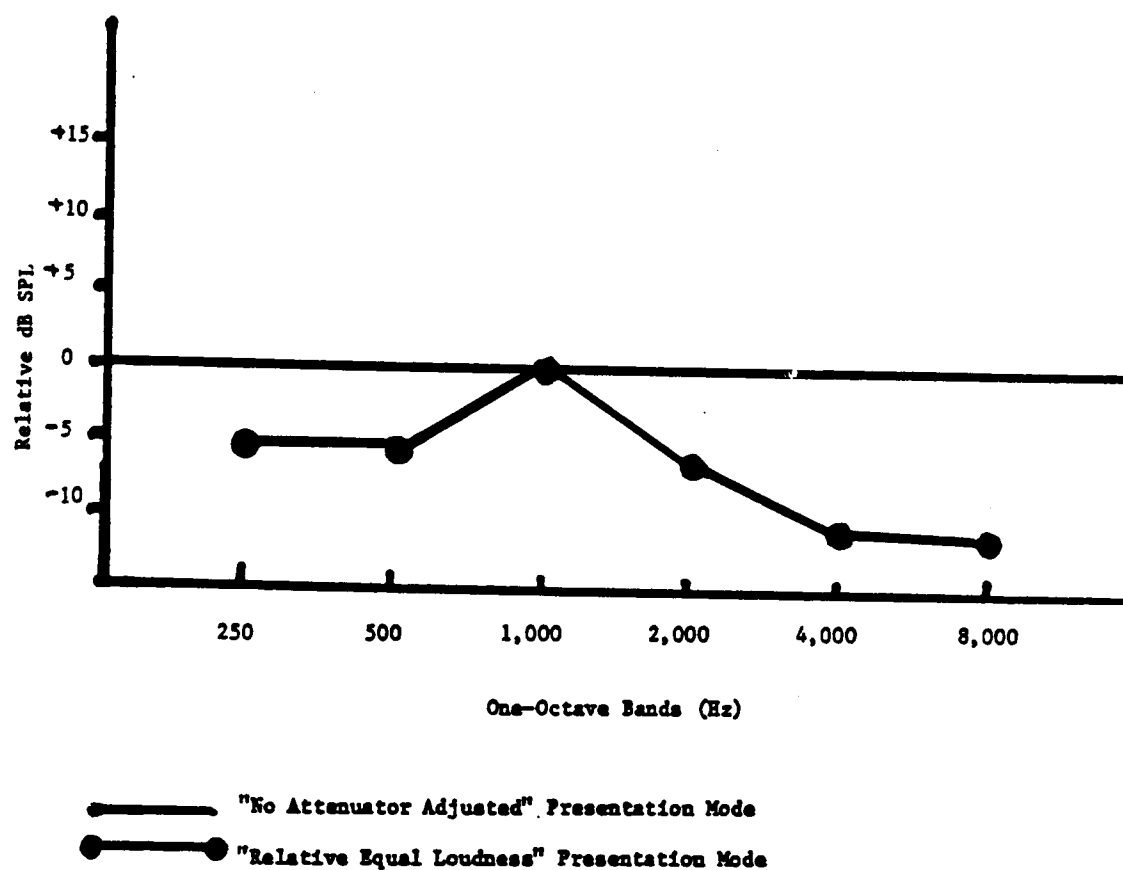
<u>Center Frequencies (Hz)</u>	<u>Cut-Off Frequencies (Hz)</u>
250	173-353
500	353-707
1,000	707-1,410
2,000	1,410-2,820
4,000	2,820-5,650
8,000	5,650-11,310

APPENDIX E

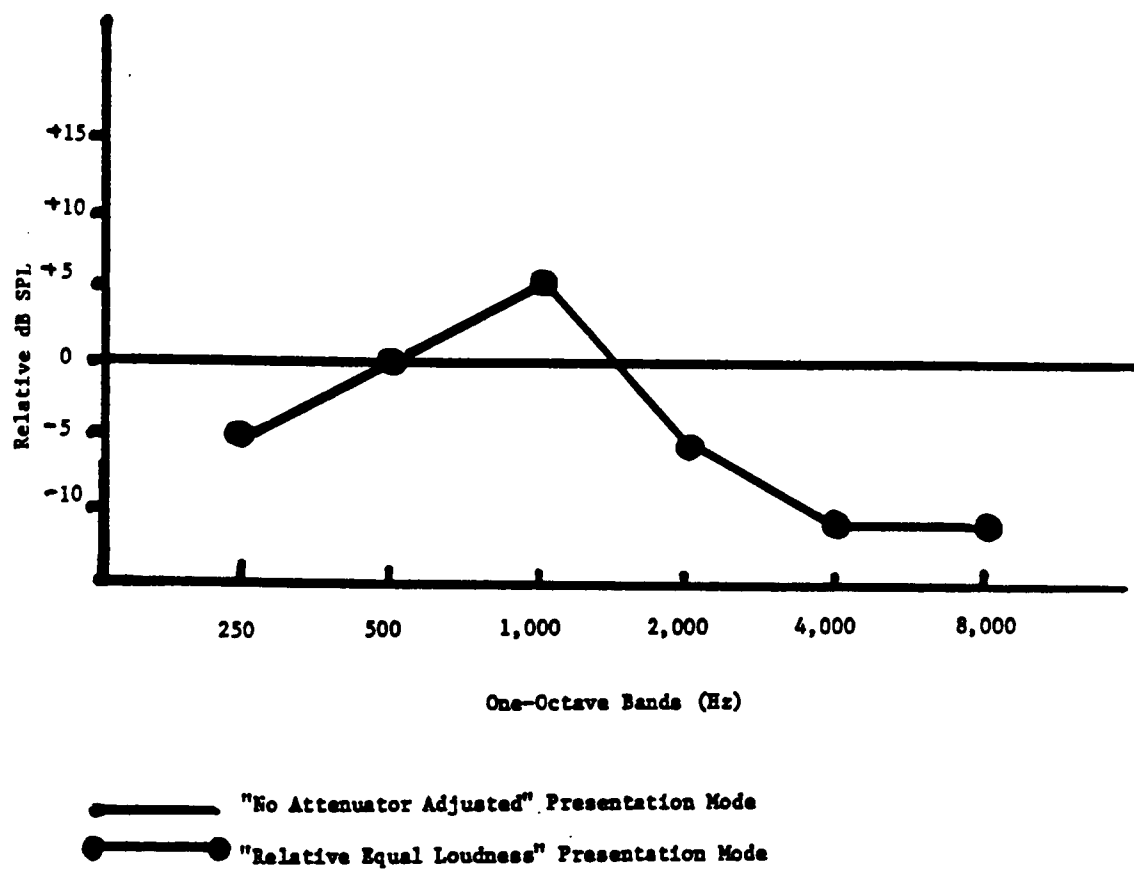
PRESENTATION MODES FOR SPEECH STIMULI



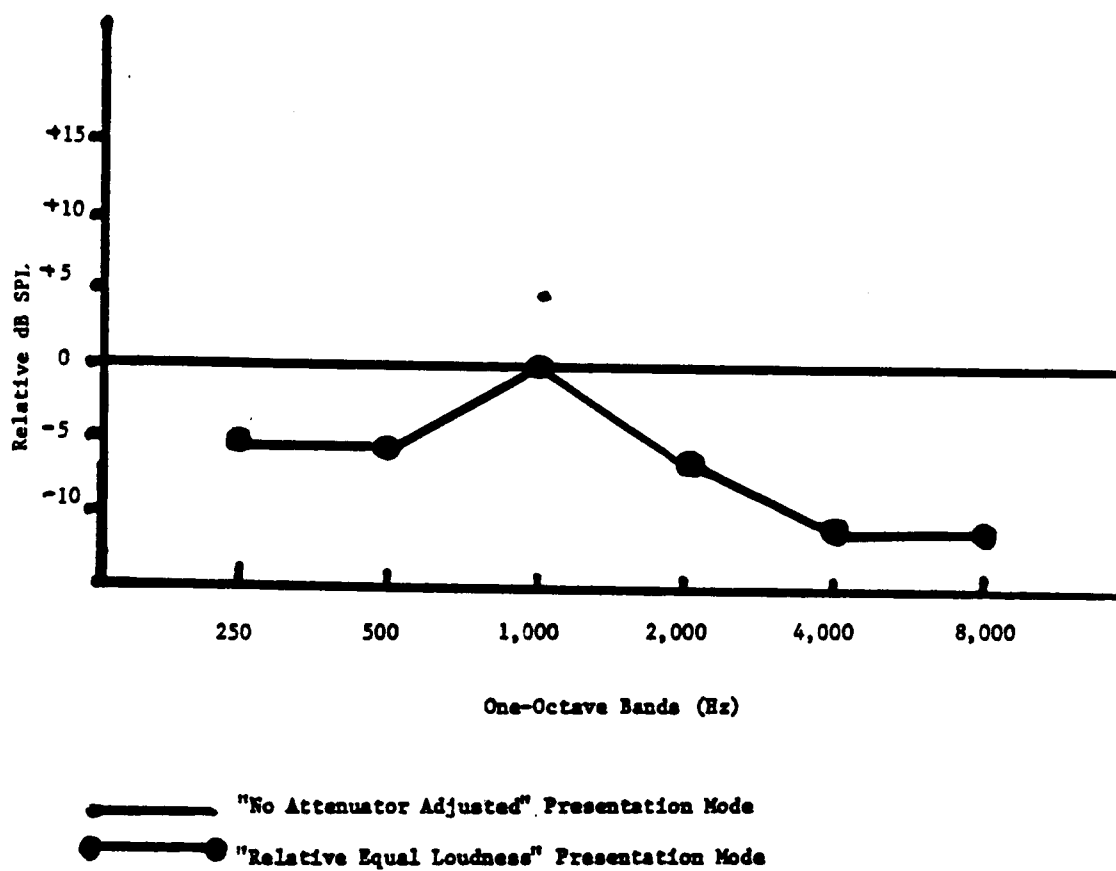
Presentation Modes for the homogeneous nonsense syllable /mumu/



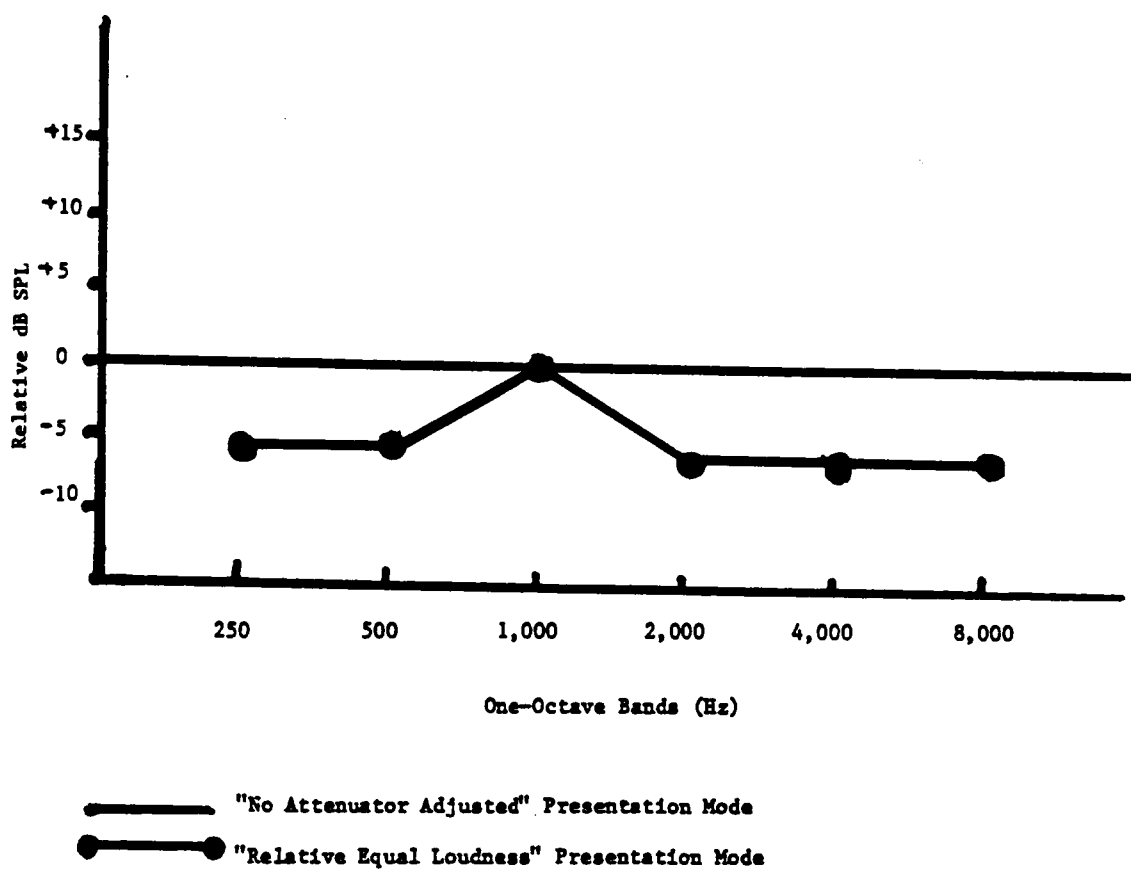
Presentation Modes for the homogeneous nonsense syllable /vovo/



Presentation modes for the homogeneous nonsense syllable /lala/



Presentation modes for the homogeneous nonsense syllable /keke/



Presentation modes for the homogeneous nonsense syllable /sisi/

APPENDIX F

EXPERIMENTAL PROTOCOL

PART I

Instructions

We will now begin Part I of the experiment. You will be involved in a task which involves listening to nonsense syllables composed of a consonant and vowel (CV- /su/), repeated twice (CVCV- /susu/), and filtered through various one-octave bands. Your task is to phonetically transcribe the CV nonsense syllable that you hear. For example, you might hear /susu/, so you will transcribe s u on the space provided for you. Note that although you will hear each CV nonsense syllable repeated twice (CVCV- /susu/), you need only transcribe s u. At times, the CV will be completely unintelligible in which case you will simply place a dash in both spaces - -. However, if the vowel is intelligible and the consonant is unintelligible, simply transcribe the vowel while putting a dash in the consonant's space. However, we encourage you to guess whenever possible. Also, even though the items on your page are numbered, they will not be numbered on the tape. Pay attention and keep up! Any questions? You will now have ten practice items:

- 1.) _ _
- 2.) _ _
- 3.) _ _
- 4.) _ _
- 5.) _ _
- 6.) _ _
- 7.) _ _
- 8.) _ _
- 9.) _ _
- 10.) _ _

Any questions? We will now begin the experiment for Part I. There will be three groups of 10 transcriptions. There will only be a brief pause between the end of one group and the beginning of the next. Please turn the page.

Form 1: Transcription for Filtered Nonsense Syllables at "No Attenuator Adjusted"

Group I

1.) _ _

2.) _ _

3.) _ _

4.) _ _

5.) _ _

6.) _ _

7.) _ _

8.) _ _

9.) _ _

10.) _ _

Group II

1.) _ _

2.) _ _

3.) _ _

4.) _ _

5.) _ _

6.) _ _

7.) _ _

8.) _ _

9.) _ _

10.) _ _

Group III

1.) _ _

2.) _ _

3.) _ _

4.) _ _

5.) _ _

6.) _ _

7.) _ _

8.) _ _

9.) _ _

10.) _ _

Form 2: Transcription for Filtered Nonsense Syllables at "Relative
Equal Loudness"Group I

1.) _ _

2.) _ _

3.) _ _

4.) _ _

5.) _ _

6.) _ _

7.) _ _

8.) _ _

9.) _ _

10.) _ _

Group II

1.) _ _

2.) _ _

3.) _ _

4.) _ _

5.) _ _

6.) _ _

7.) _ _

8.) _ _

9.) _ _

10.) _ _

Group III

1.) _ _

2.) _ _

3.) _ _

4.) _ _

5.) _ _

6.) _ _

7.) _ _

8.) _ _

9.) _ _

10.) _ _

PART II

Instructions

Now you will complete a task which requires you to judge the perceived pitch of unfiltered CV nonsense syllables repeated twice. You will be hearing the nonsense syllables in pairs. Your task is to circle which member of the pair is higher in pitch. Remember, circle the member of the pair which is higher in pitch. Also, even though the items on your response sheet are numbered, they are not numbered on the tape. We will now have ten practice items. Any questions? Remember, higher in pitch!

1.) 1 2

2.) 1 2

3.) 1 2

4.) 1 2

5.) 1 2

6.) 1 2

7.) 1 2

8.) 1 2

9.) 1 2

10.) 1 2

Any questions? We will now begin the experiment. Please turn the page.

Remember, higher in pitch!

- | | |
|----------|----------|
| 1.) 1 2 | 11.) 1 2 |
| 2.) 1 2 | 12.) 1 2 |
| 3.) 1 2 | 13.) 1 2 |
| 4.) 1 2 | 14.) 1 2 |
| 5.) 1 2 | 15.) 1 2 |
| 6.) 1 2 | 16.) 1 2 |
| 7.) 1 2 | 17.) 1 2 |
| 8.) 1 2 | 18.) 1 2 |
| 9.) 1 2 | 19.) 1 2 |
| 10.) 1 2 | 20.) 1 2 |

PART III

Instructions

We will now begin Part III of the experiment. You will be rating the quality and the essence of the filtered nonsense syllables on a scale of 1 (lowest quality) to 5 (highest quality). In your ratings, you will know what nonsense syllables you are rating as they will be phonetically transcribed for each item. We will now play three examples of how to use the number scale:

For the highest quality, this would be an example of a 5 rating:

Quality					
low			high		
/susu/	1	2	3	4	5

For the lowest quality, this would be an example of a 1 rating:

Quality					
low			high		
/susu/	1	2	3	4	5

For average quality, this would be an example of a 3 rating:

Quality					
low			high		
/susu/	1	2	3	4	5

Any questions? You will now have ten practice items. Try and concentrate on the quality and essence of the nonsense syllable while trying to ignore how the filter itself sounds. For example, a nonsense syllable may sound "different" than natural production when filtered through a certain band. Do not use "different than normal" as a criteria for "high quality," but rather "high quality" may imply more enhanced than normal or "better" than normal production.

Quality					
low			high		
1.) lolo	1	2	3	4	5

	low			high	
2.) susu	1	2	3	4	5
3.) kiki	1	2	3	4	5
4.) mama	1	2	3	4	5
5.) sisi	1	2	3	4	5
6.) vovo	1	2	3	4	5
7.) lili	1	2	3	4	5
8.) mumu	1	2	3	4	5
9.) sasa	1	2	3	4	5
10.) vuvu	1	2	3	4	5

Any questions? We will now begin Part III of the experiment. As in Part I, the ratings will come in three groups of ten. This time though, each group of 10 will be on a different page. When you finish with one group of ten, quickly turn to the next page to begin rating the next group. Turn to the next page.

Group I

	Quality				
	low				high
1.) /mumu/	1	2	3	4	5
2.) /sisi/	1	2	3	4	5
3.) /vovo/	1	2	3	4	5
4.) /keke/	1	2	3	4	5
5.) /sisi/	1	2	3	4	5
6.) /laɭa/	1	2	3	4	5
7.) /mumu/	1	2	3	4	5
8.) /vovo/	1	2	3	4	5
9.) /laɭa/	1	2	3	4	5
10.) /keke/	1	2	3	4	5

Group II

	Quality				
	low				high
1.) /keke/	1	2	3	4	5
2.) /mumu/	1	2	3	4	5
3.) /la la/	1	2	3	4	5
4.) /vovo/	1	2	3	4	5
5.) /keke/	1	2	3	4	5
6.) /sisi/	1	2	3	4	5
7.) /la la/	1	2	3	4	5
8.) /sisi/	1	2	3	4	5
9.) /vovo/	1	2	3	4	5
10.) /mumu/	1	2	3	4	5

Group III

	Quality				
	low				high
1.) /lɑlɑ/	1	2	3	4	5
2.) /mumu/	1	2	3	4	5
3.) /vovo/	1	2	3	4	5
4.) /lɑlɑ/	1	2	3	4	5
5.) /keke/	1	2	3	4	5
6.) /sisi/	1	2	3	4	5
7.) /vovo/	1	2	3	4	5
8.) /sisi/	1	2	3	4	5
9.) /keke/	1	2	3	4	5
10.) /mumu/	1	2	3	4	5

POST-EXPERIMENT QUESTIONNAIRE

Thank you for participating in this study. All of your responses and identity will be kept anonymous. However, it would be appreciated if you could answer the following questions:

Sex M F

Age _____

Occupation _____

Years experience _____

Please circle the appropriate answer:

- | | | |
|-------------------------------------|-----|----|
| 1.) Was the task difficult? | Yes | No |
| 2.) Was the signal clear? | Yes | No |
| 3.) Were there any distractions? | Yes | No |
| 4.) Were you motivated? | Yes | No |
| 5.) Were you sure of your responses | Yes | No |

How could this task have been improved to maximize your concentration and performance?

THANK YOU!

APPENDIX G

OCTAVE BAND MEASURES OF AMBIENT NOISE IN SOUND-TREATED SUITE
AND ALLOWABLE LEVELS FOR AUDIOMETRIC TEST BOOTHS

OCTAVE BAND MEASUREMENTS OF SOUND-TREATED SUITE

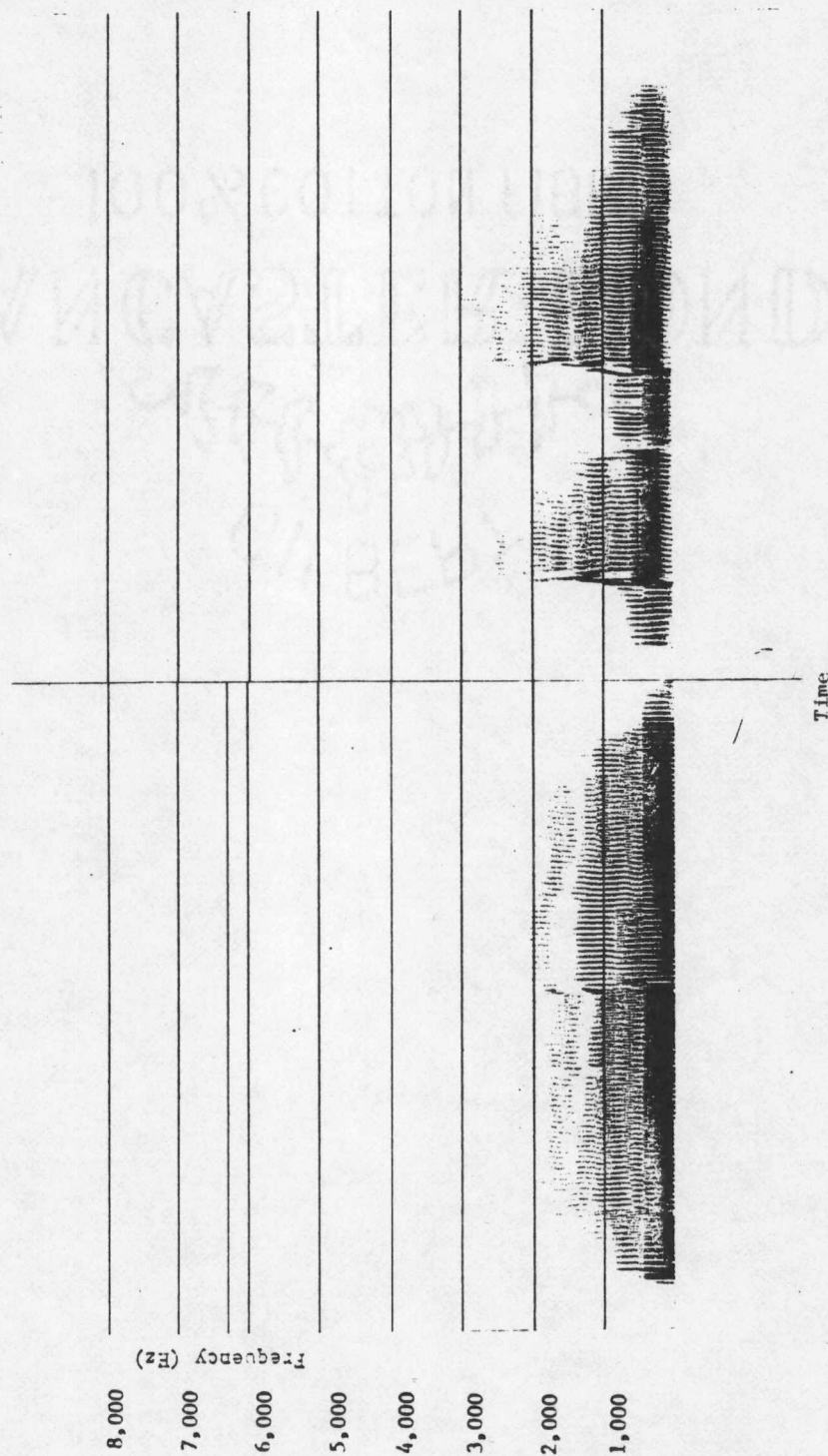
<u>Center Frequency (Hz)</u>	<u>Level dB SPL</u>
62	40
125	20.5
250	12
500	8
1,000	7
2,000	8
4,000	8.5
8,000	7.5

CRITERIA FOR BACKGROUND NOISE IN AUDIOMETER BOOTHS

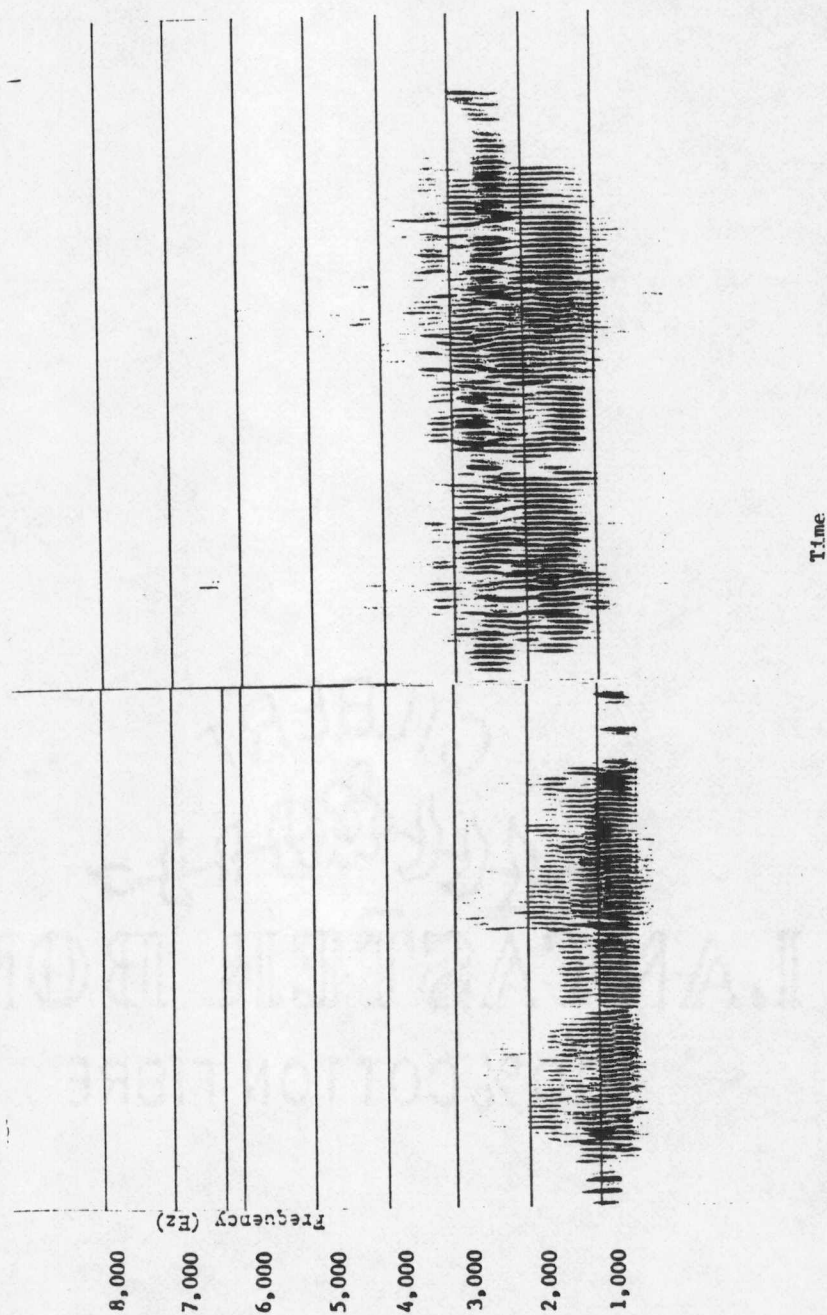
<u>Center Frequency (Hz)</u>	<u>Allowable Level dB SPL</u>
125	34.5
250	23
500	21.5
1,000	29.5
2,000	35.4
4,000	42
8,000	45

APPENDIX H

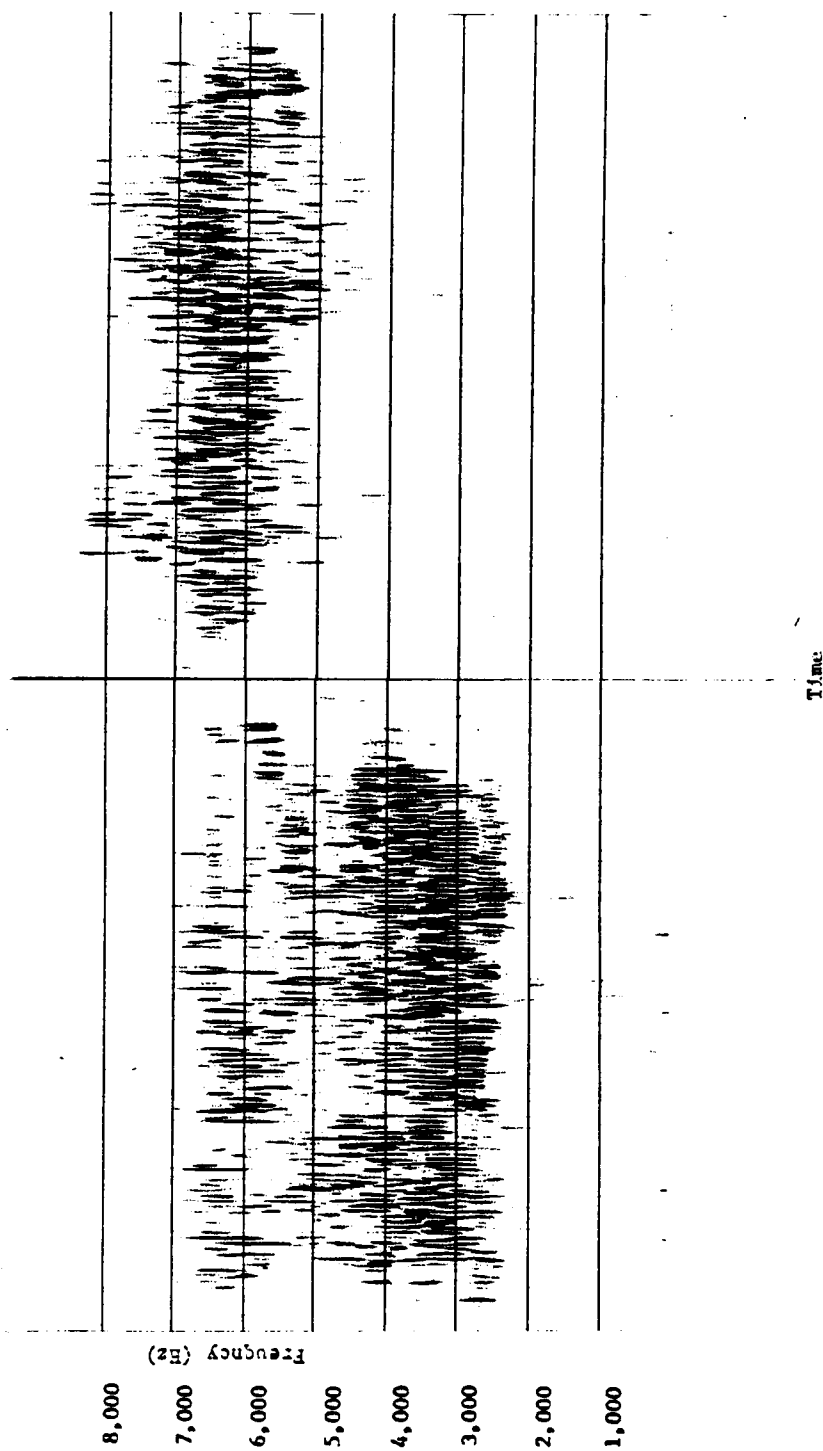
SPECTROGRAPHIC ANALYSES OF SPEECH STIMULI USED IN THIS STUDY



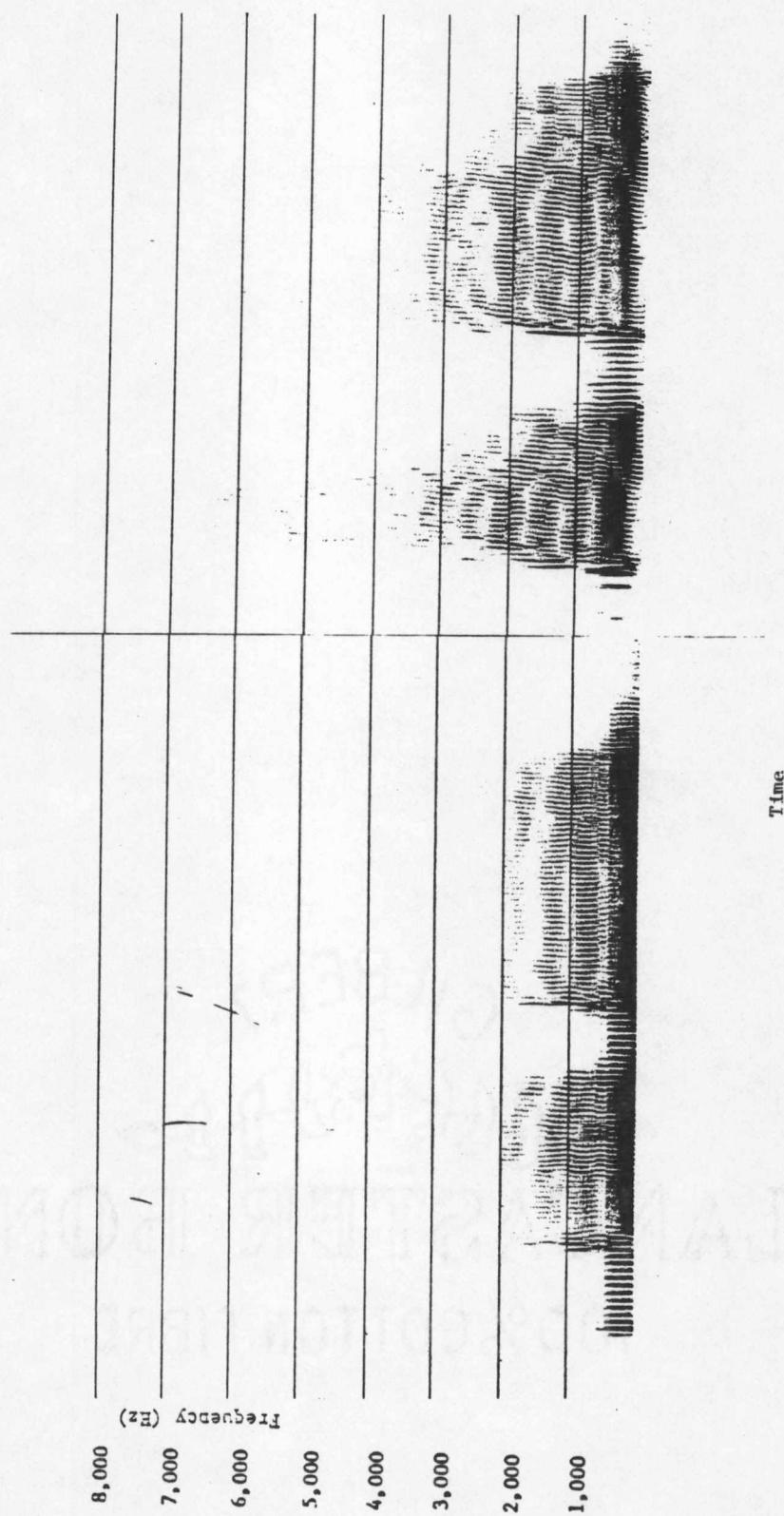
Wide-band analysis of the homogeneous nonsense syllable /mumu/ filtered through a 1,000 and 2,000 Hz one-octave filter band.



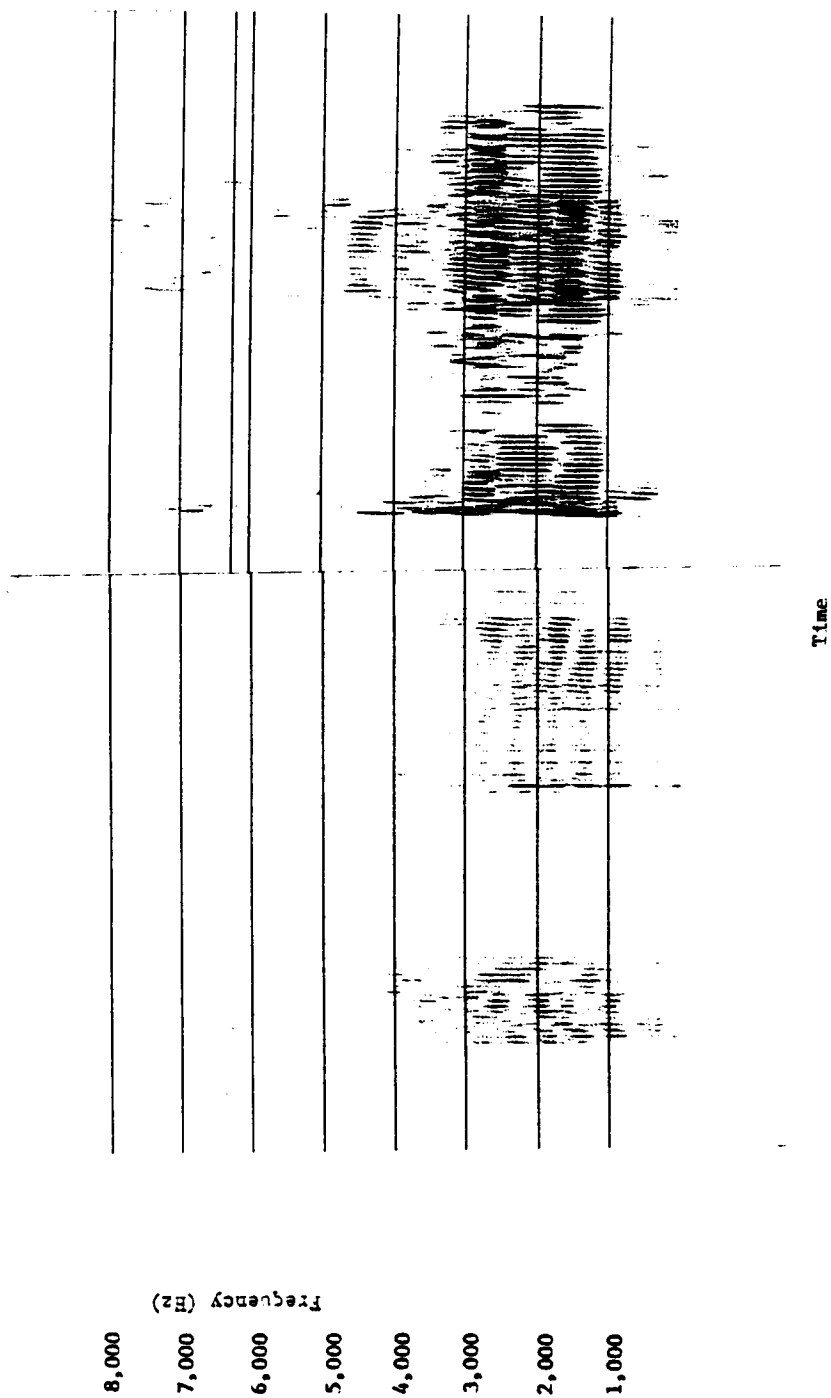
Wide-band analysis of the homogeneous nonsense syllable /mumu/ filtered through a 250 and 500 Hz one-octave filter band.



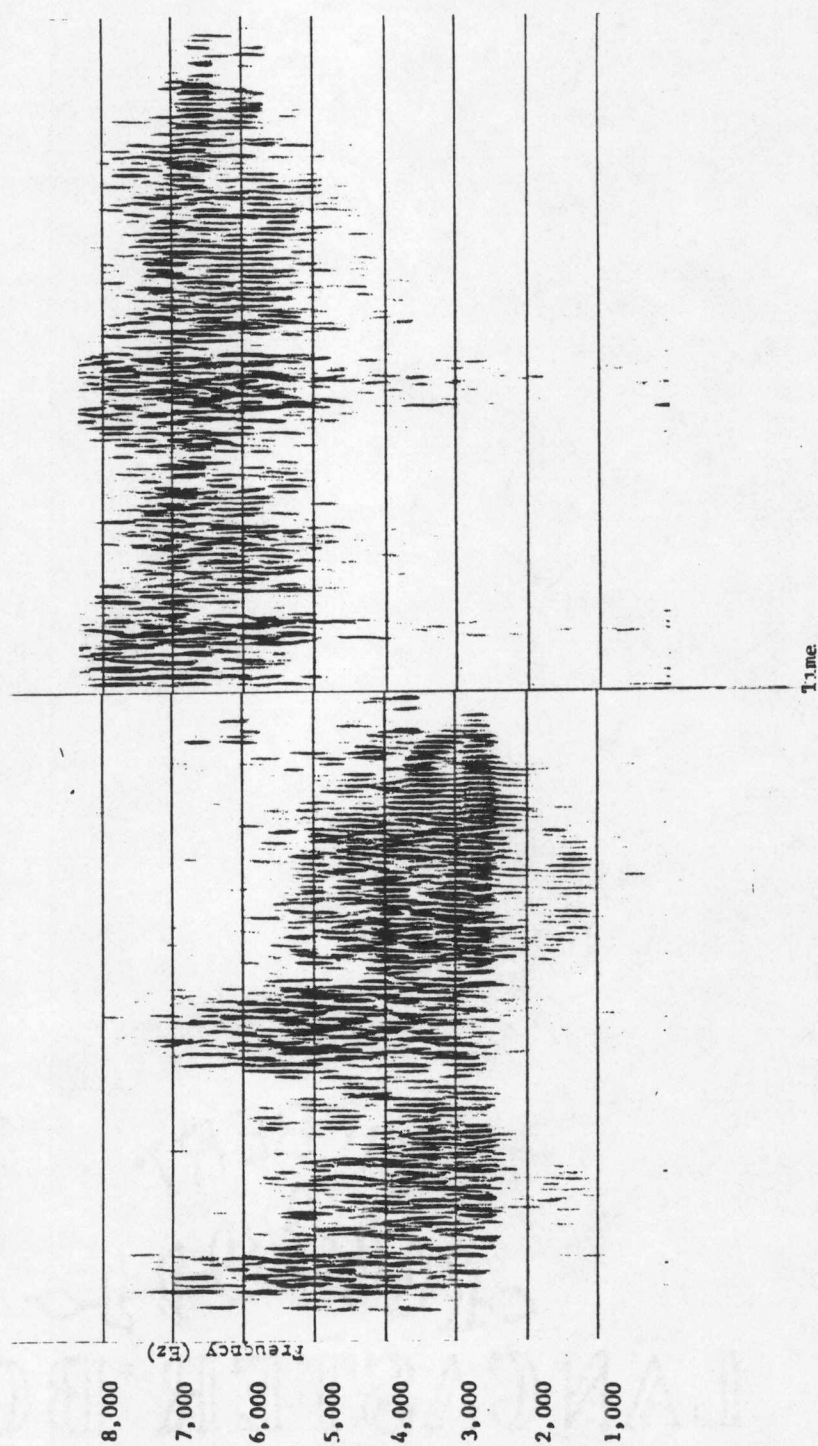
Wide-band analysis of the homogeneous nonsense syllable /mumu/ filtered through a 4,000 and 8,000 Hz one-octave filter band.



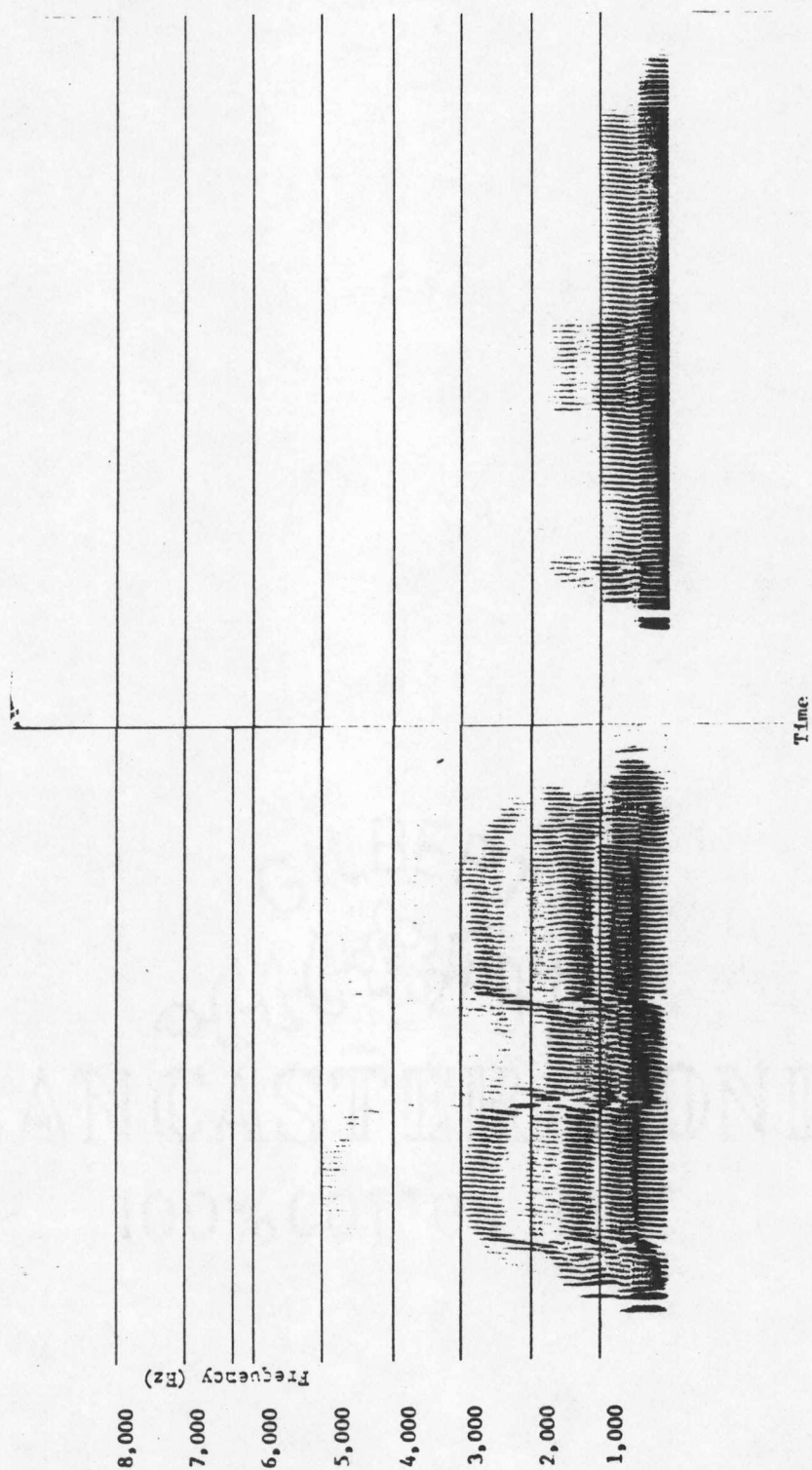
Wide-band analysis of the homogeneous nonsense syllable /vovo/ filtered through a 250 and 500 Hz one-octave filter band.



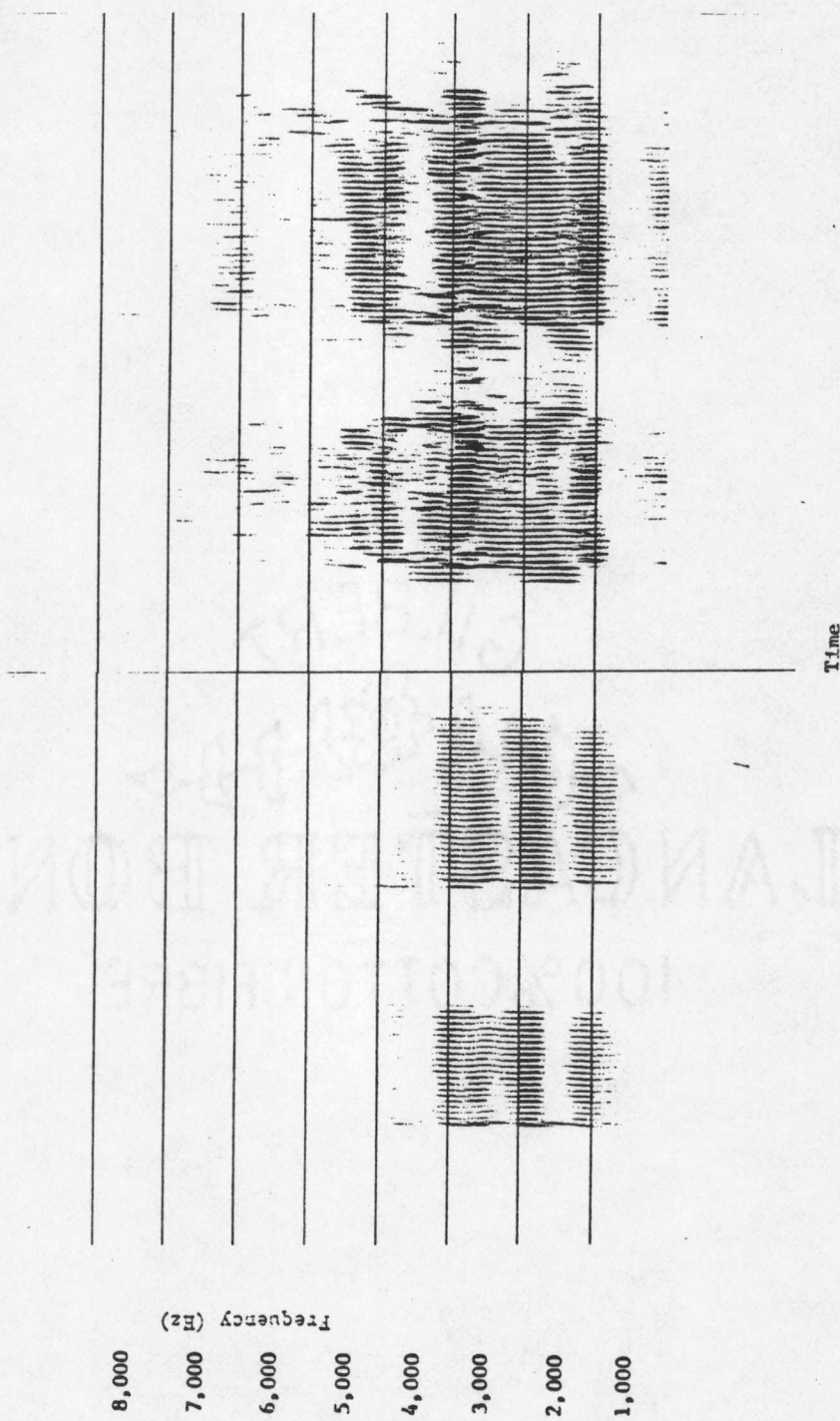
Wide-band analysis of the homogeneous nonsense syllable /vovo/ filtered through a 1,000 and 2,000 Hz one-octave filter band.



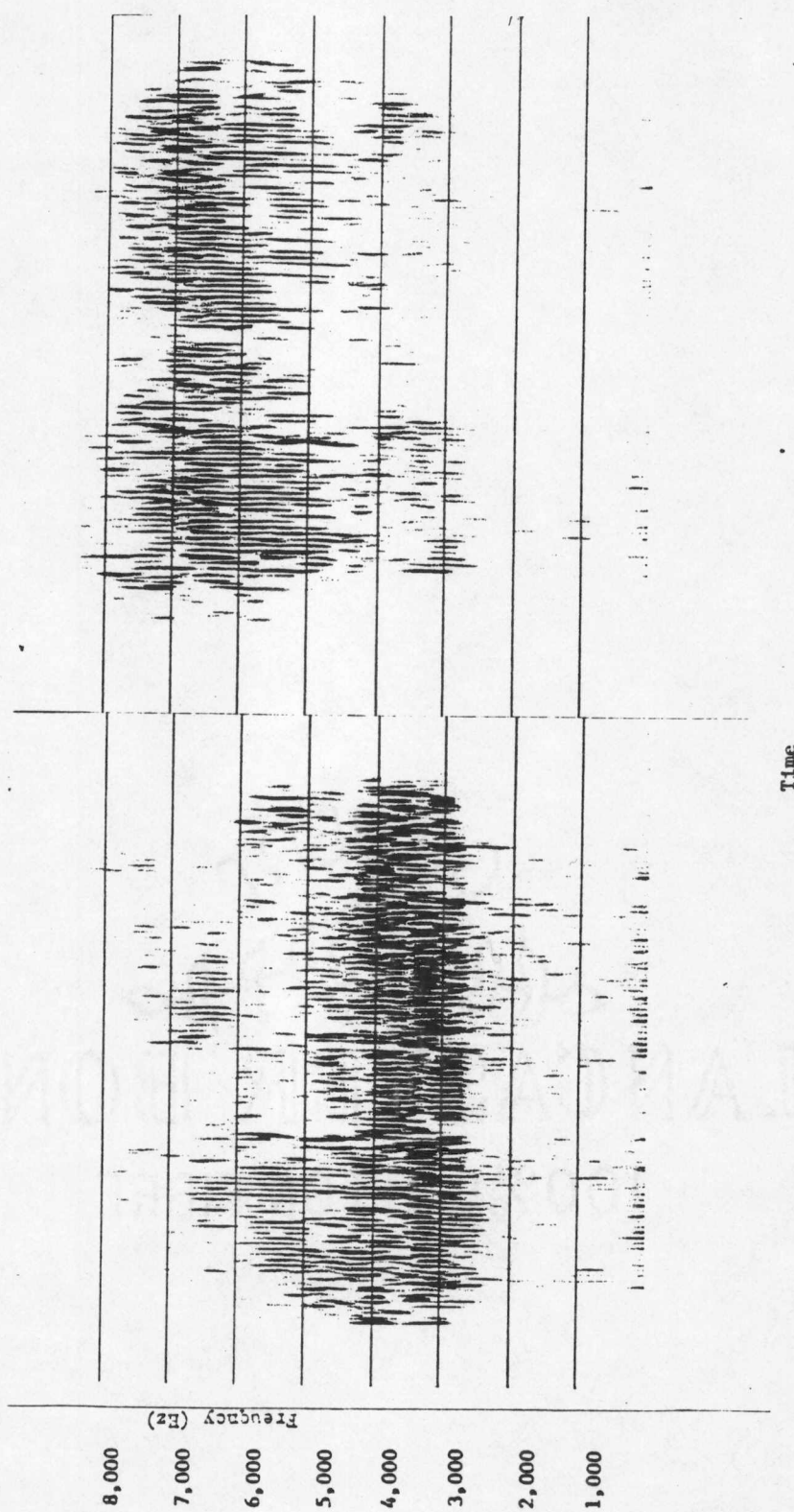
Wide-band analysis of the homogeneous nonsense syllable /vovo/ filtered through a 4,000 and 8,000 Hz one-octave filter band.



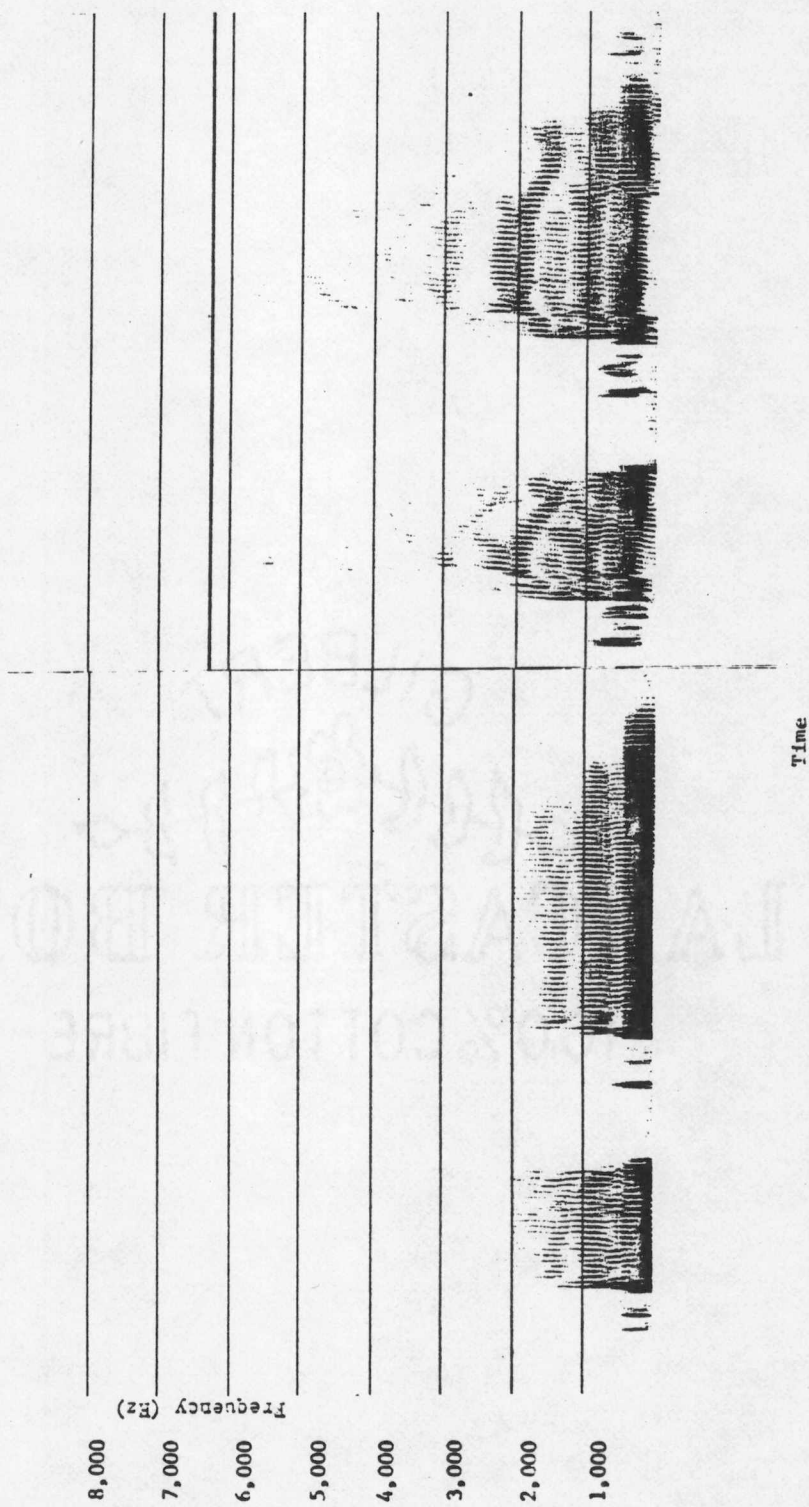
Wide-band analysis of the homogeneous nonsense syllable /laja/ filtered through a 250 and 500 Hz one-octave filter band.



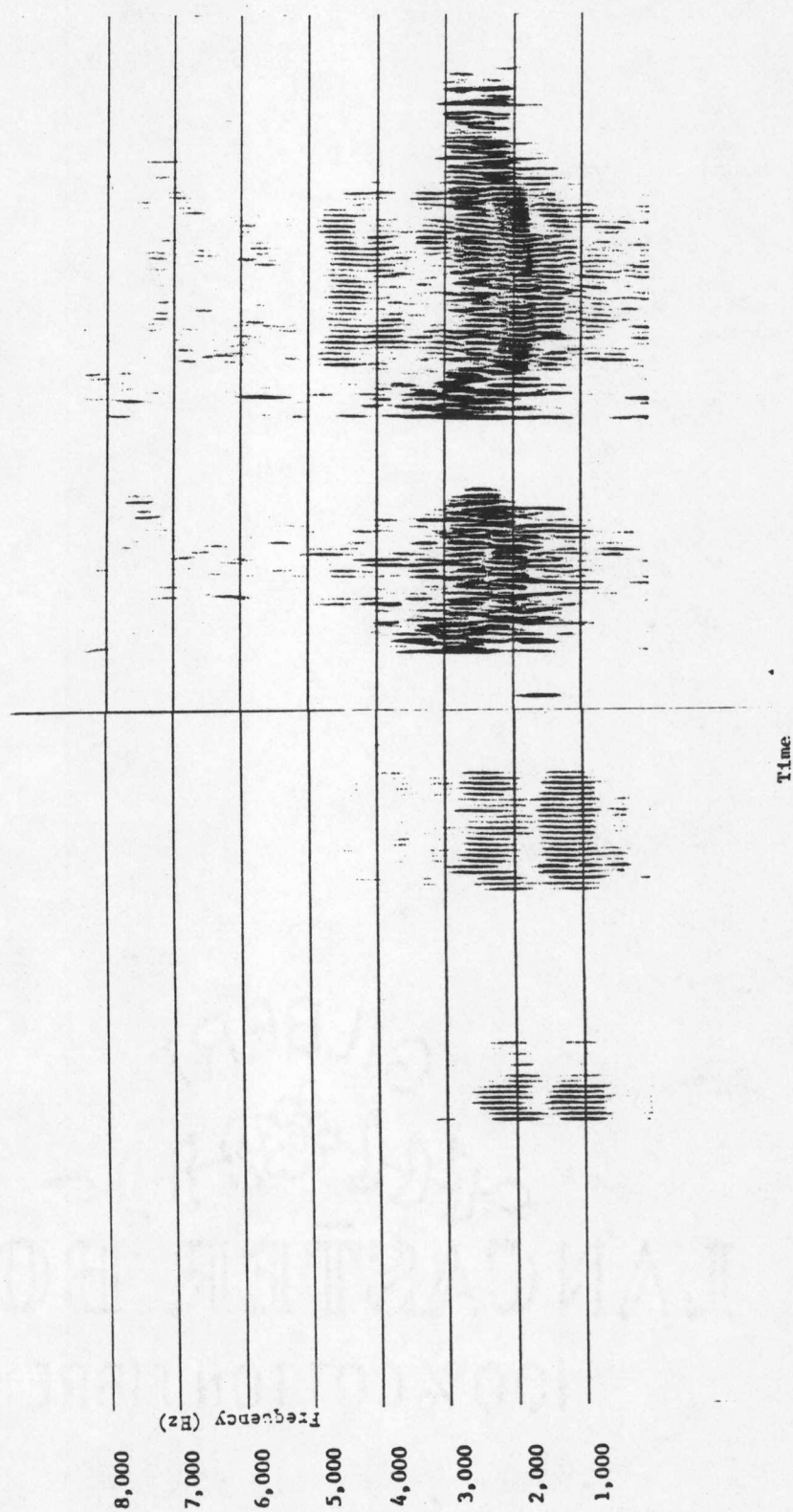
Wide-band analysis of the homogeneous nonsense syllable /la₁la/ filtered through a 1,000 and 2,000 Hz one-octave filter band.



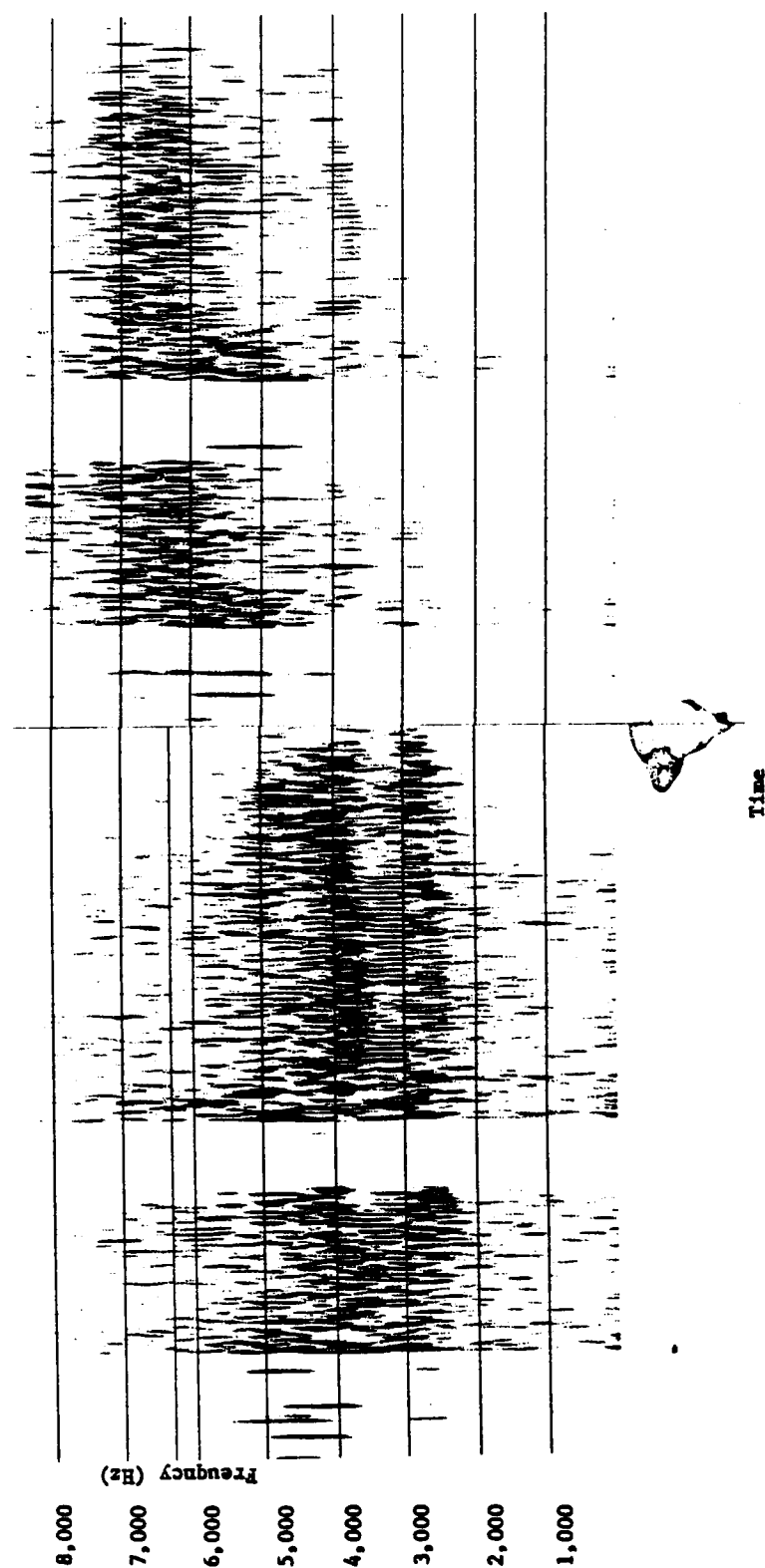
Wide-band analysis of the homogeneous nonsense syllable /laɪa/ filtered through a 4,000 and 8,000 Hz one-octave filter band.



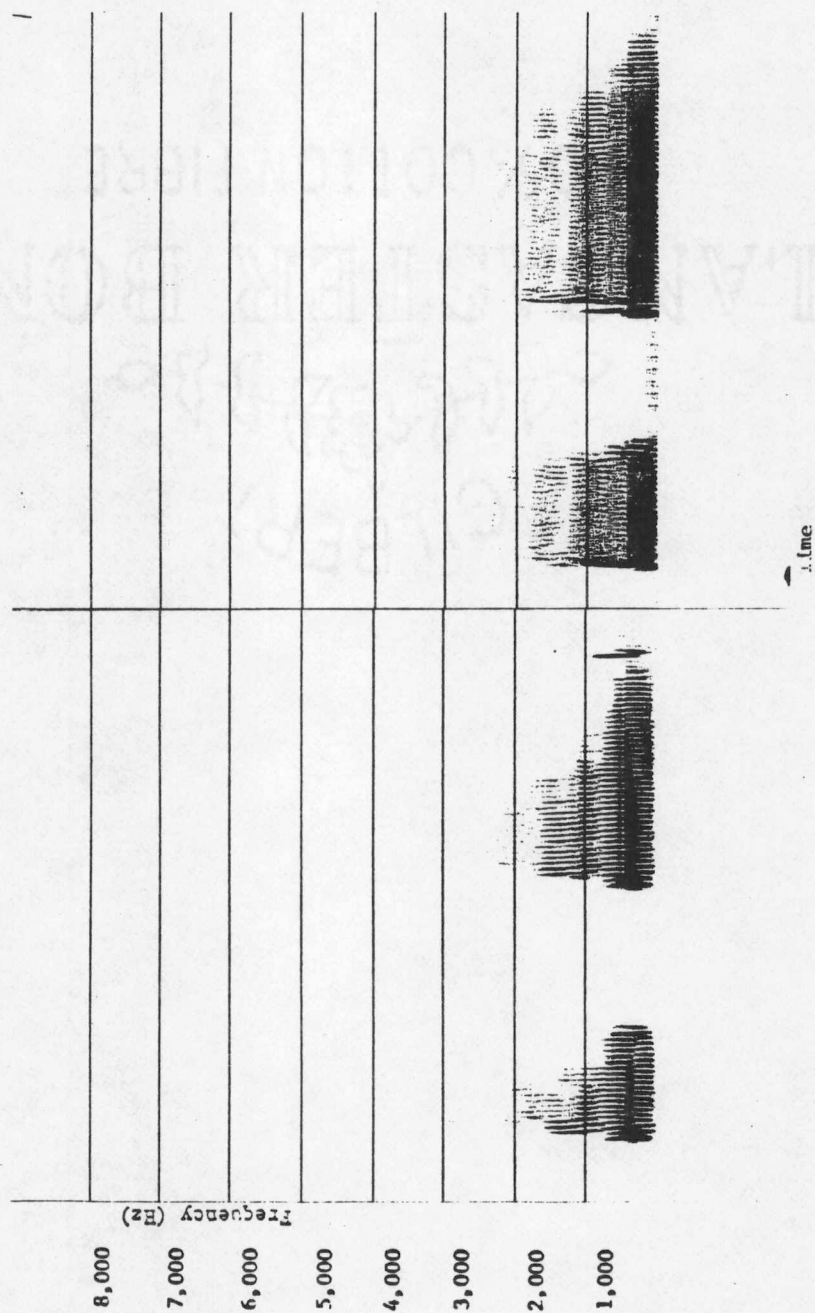
Wide-band analysis of the homogeneous nonsense syllable /keke/ filtered through a 250 and 500 Hz one-octave filter band.



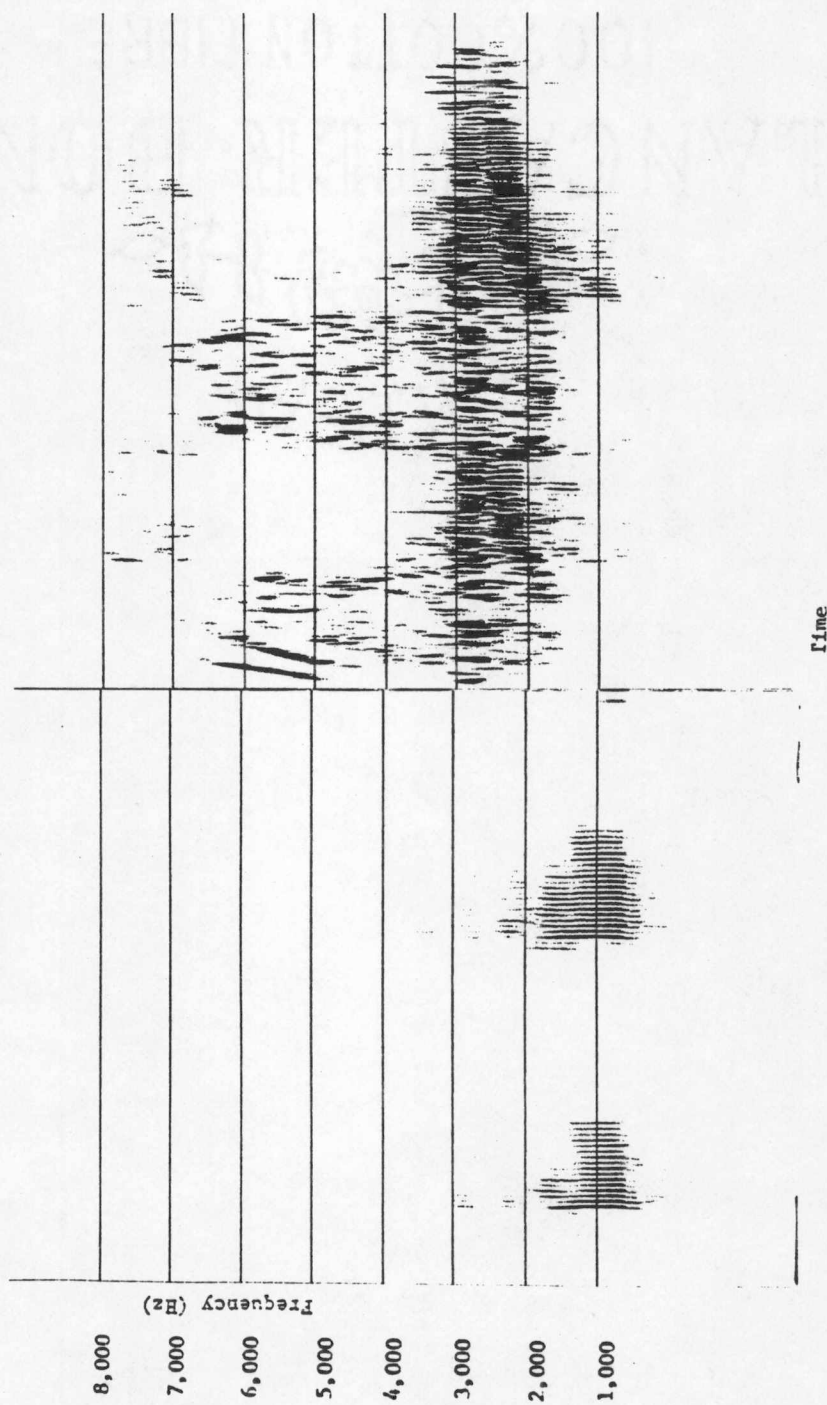
Wide-band analysis of the homogeneous nonsense syllable /keke/ filtered through a 1,000 and 2,000 Hz one-octave filter band.



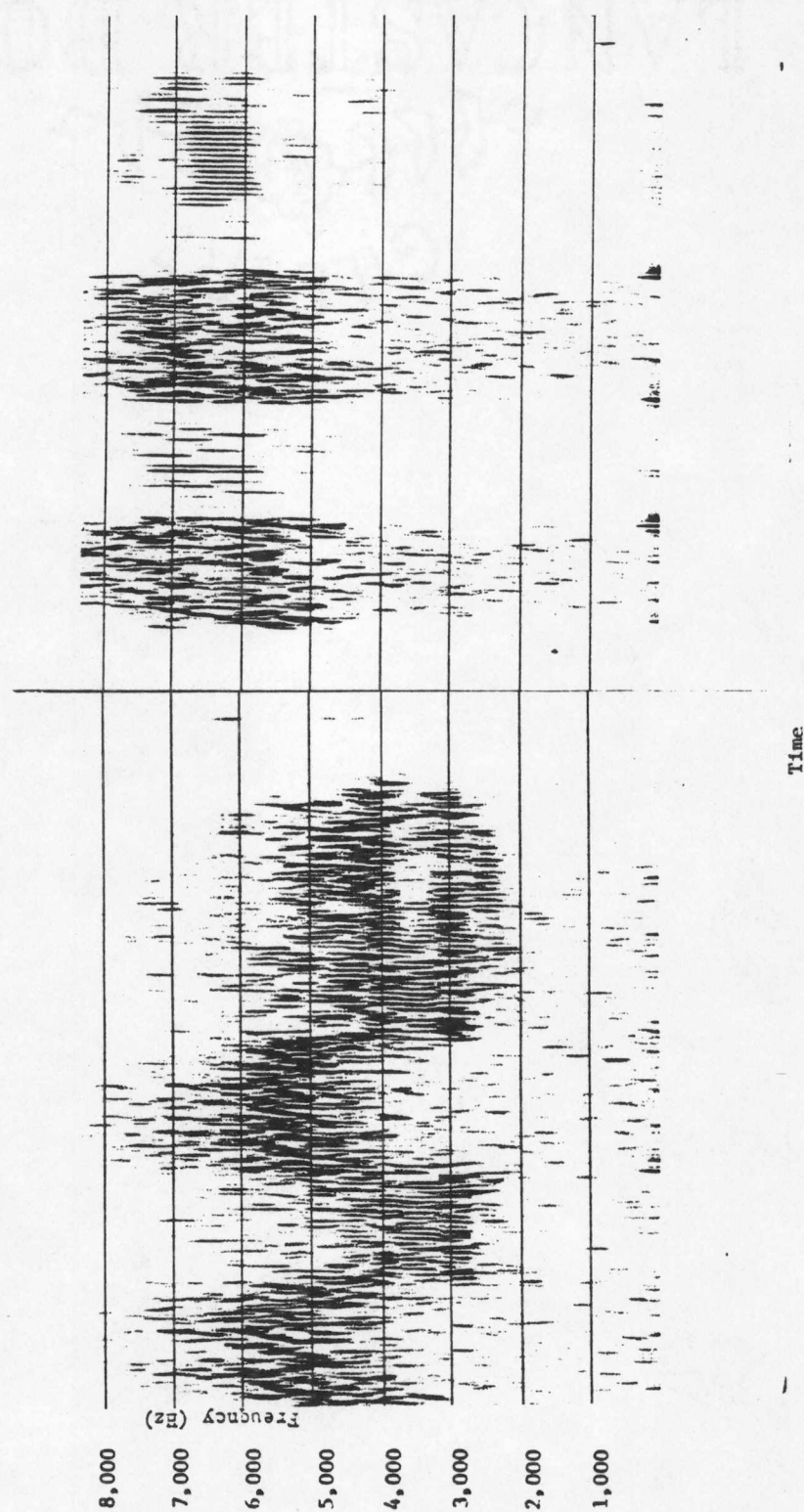
Wide-band analysis of the homogeneous nonsense syllable /keke/ filtered through a 4,000 and 8,000 Hz one-octave filter band.



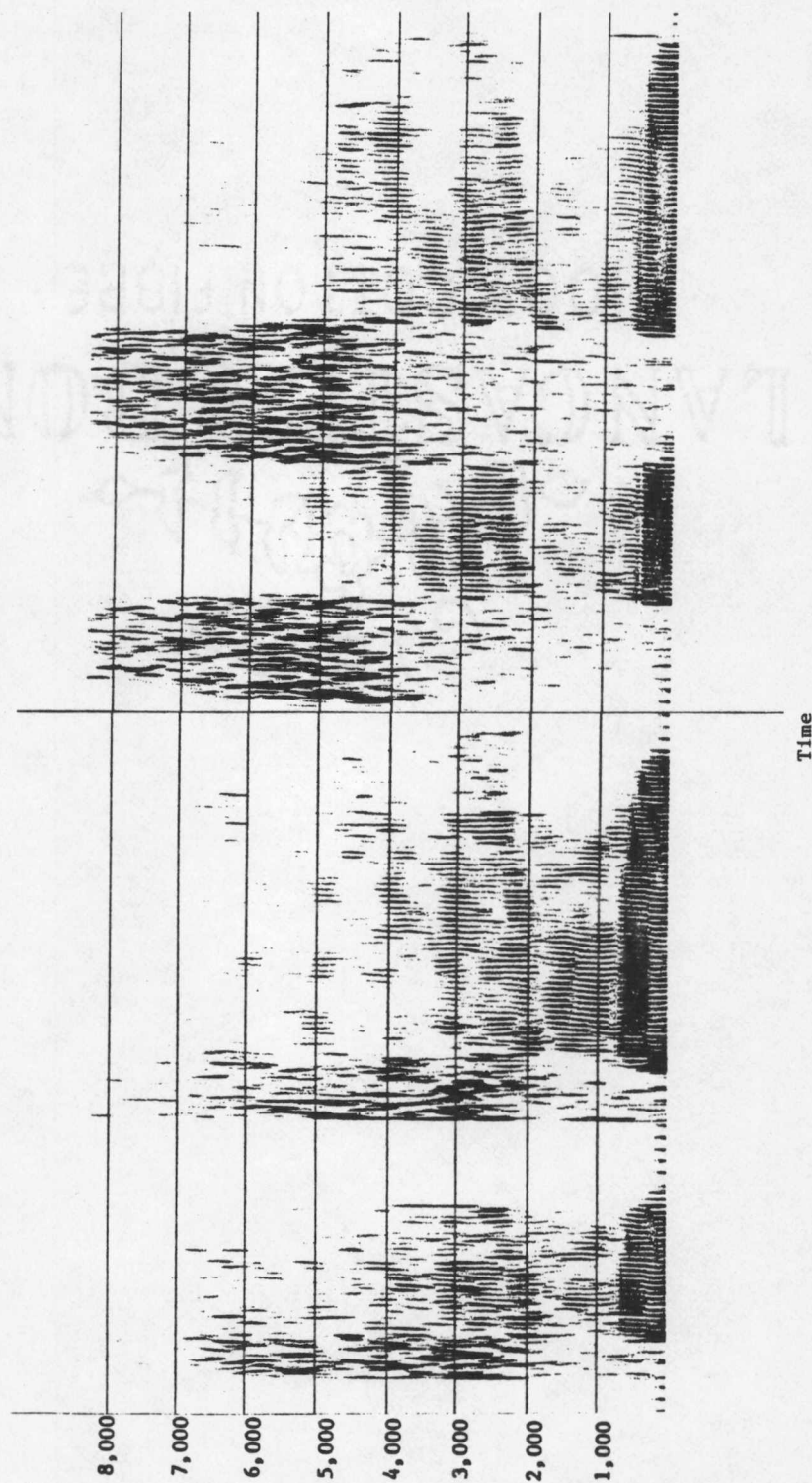
Wide-band analysis of the homogeneous nonsense syllable /sisi/ filtered through a 250 and 500 Hz one-octave filter band.



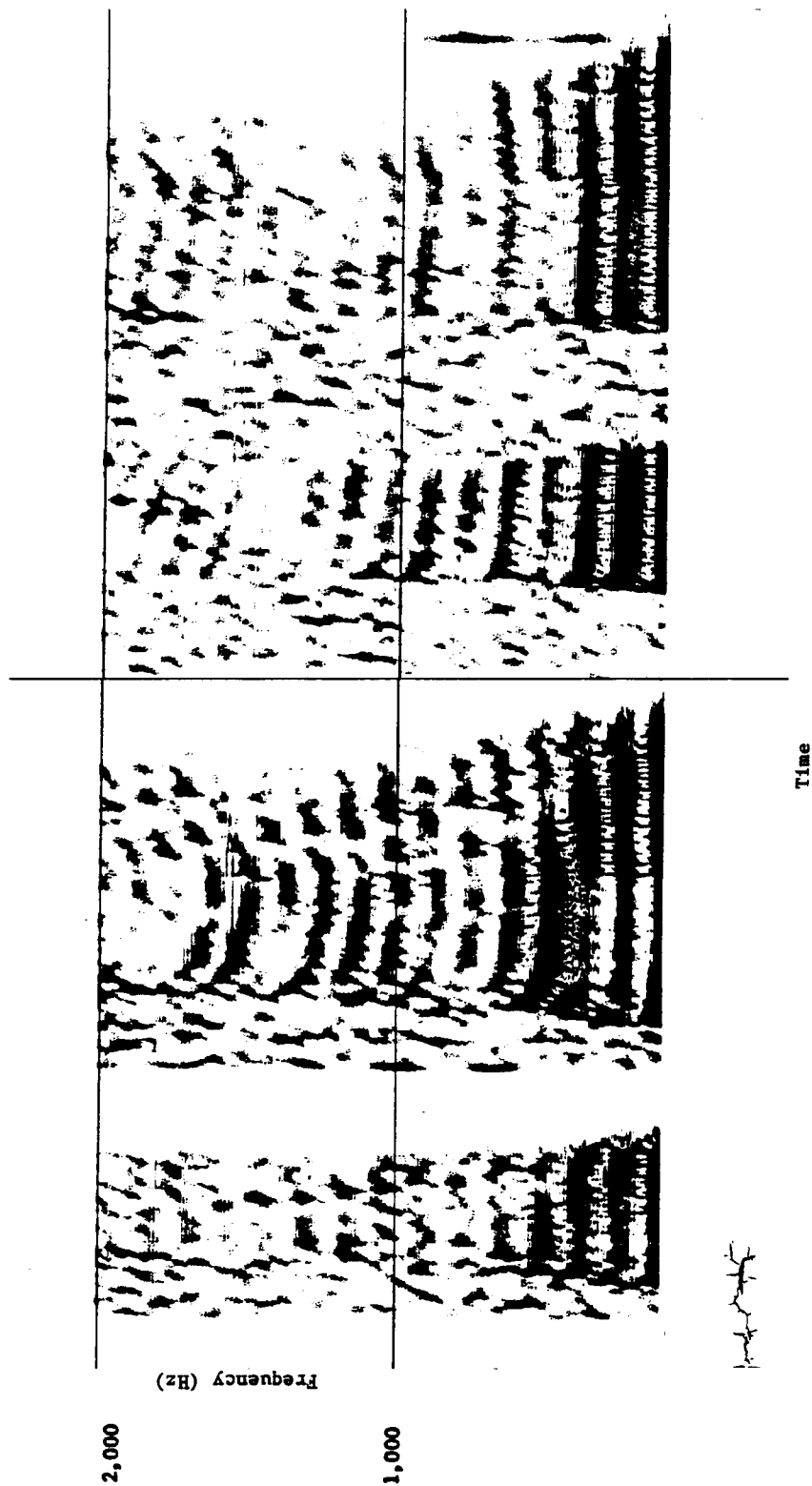
Wide-band analysis of the homogeneous nonsense syllable /sisi/ filtered through a 1,000 and 2,000 Hz one-octave filter band.



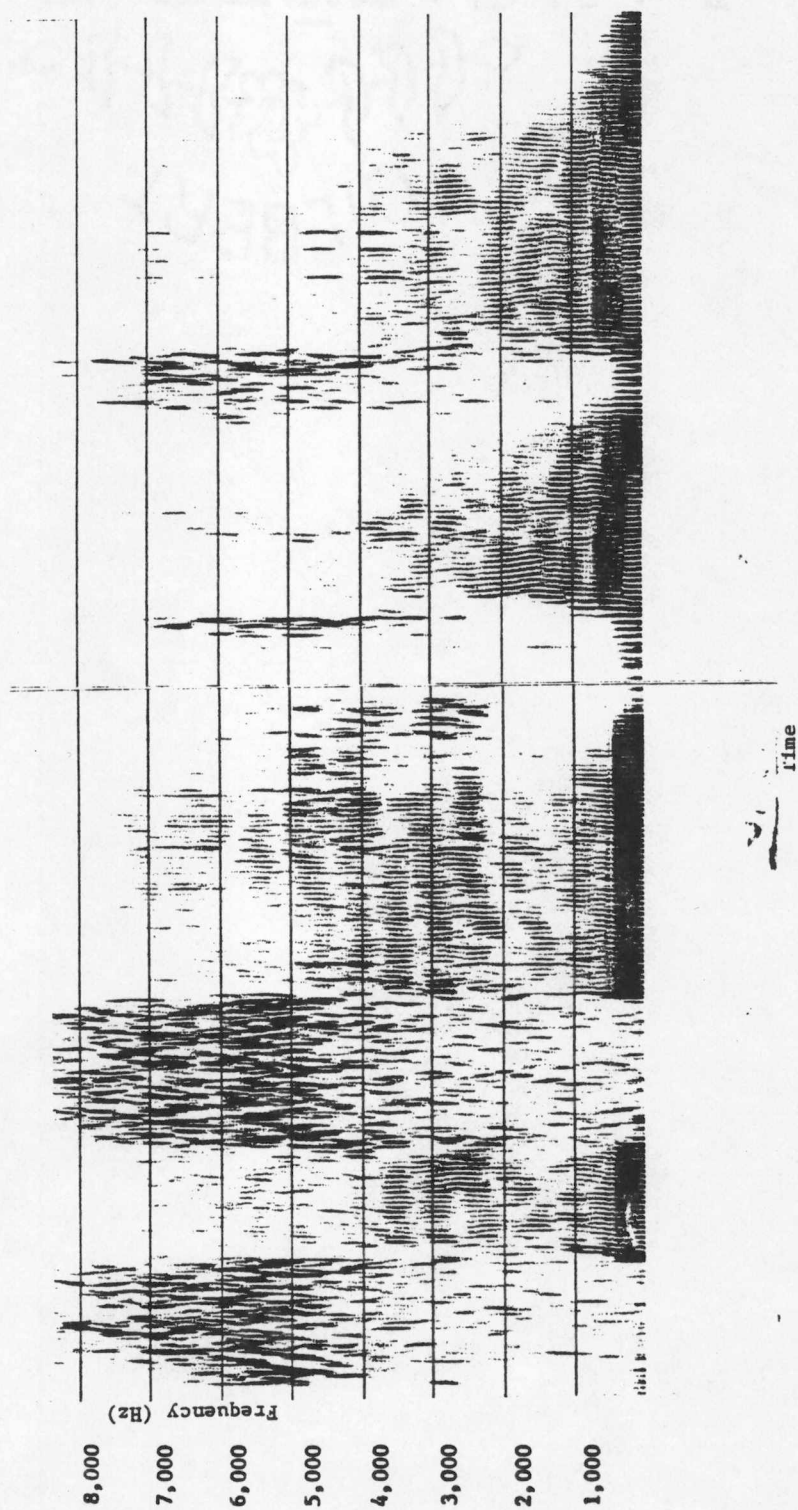
Wide-band analysis of the homogeneous nonsense syllable /sisi/ filtered through a 4,000 and 8,000 Hz one-octave filter band.



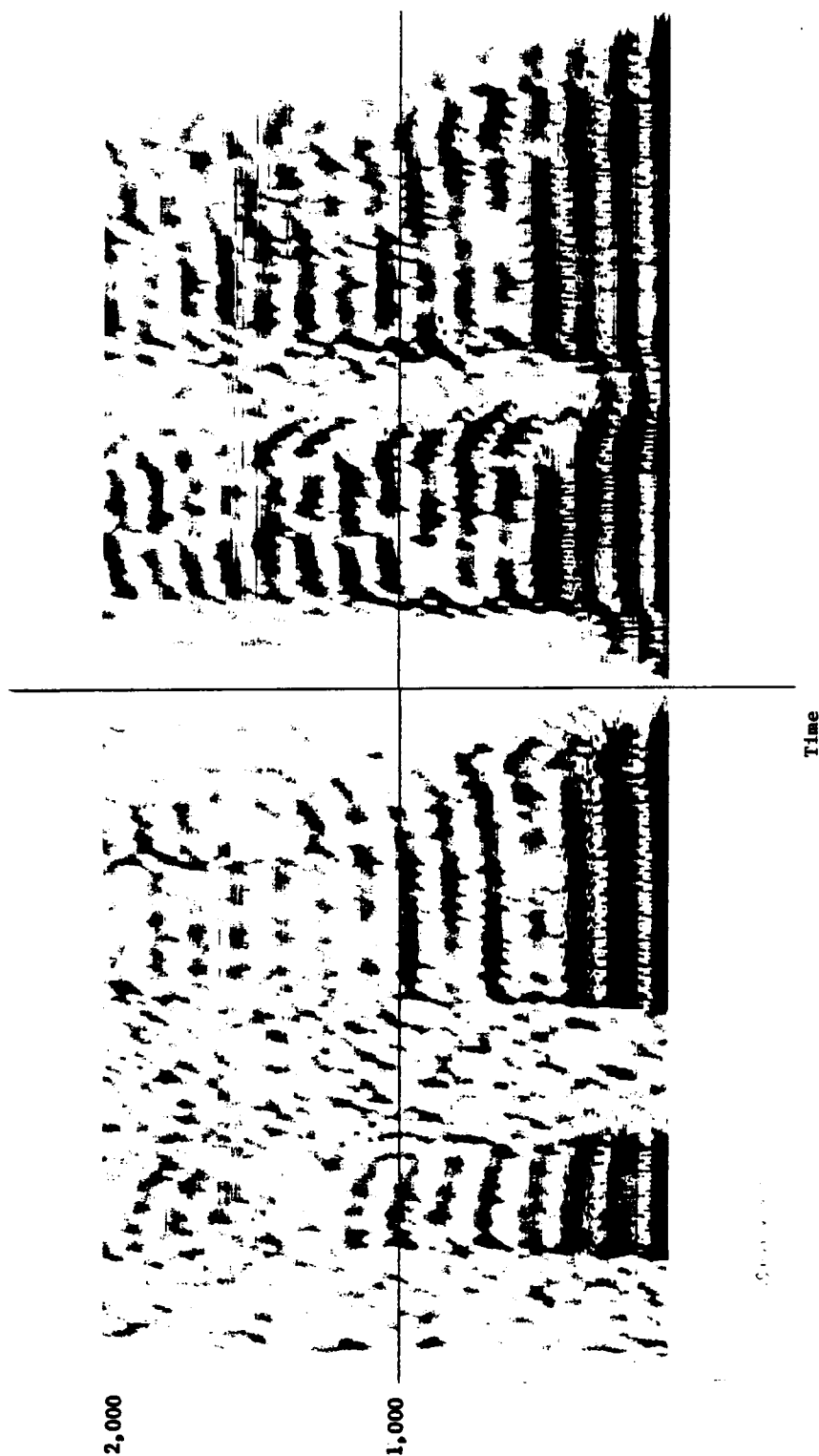
Wide-band analysis of the homogeneous nonsense syllables /keke/ - /sisi/ to 8,000 Hz.



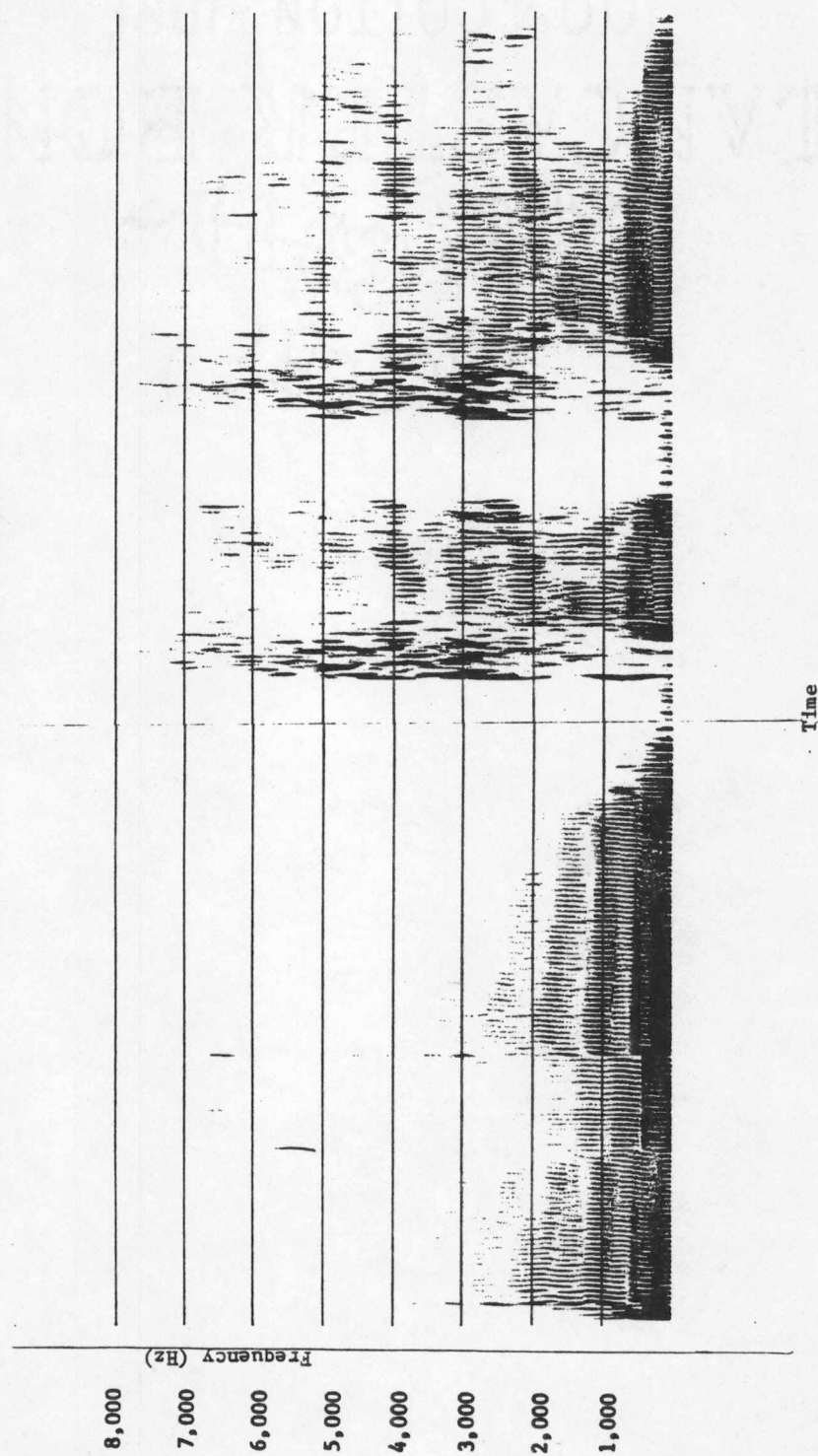
Narrow-band analysis of the homogeneous nonsense syllables /keke/ - /sisi/ expanded to 2,000 Hz.



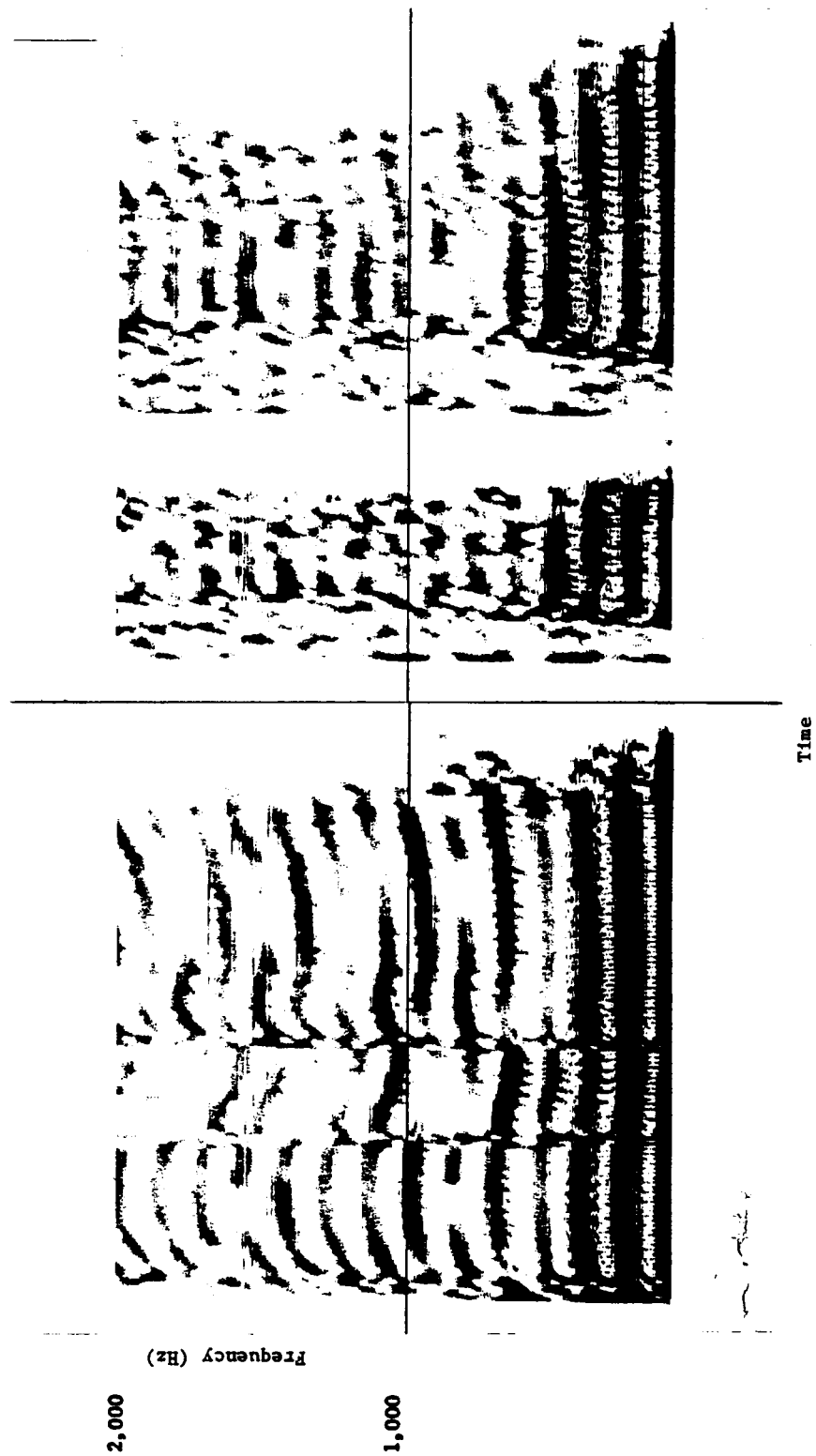
Wide-band analysis of the homogeneous nonsense syllables /sisi/ - /vovo/ to 8,000 Hz.



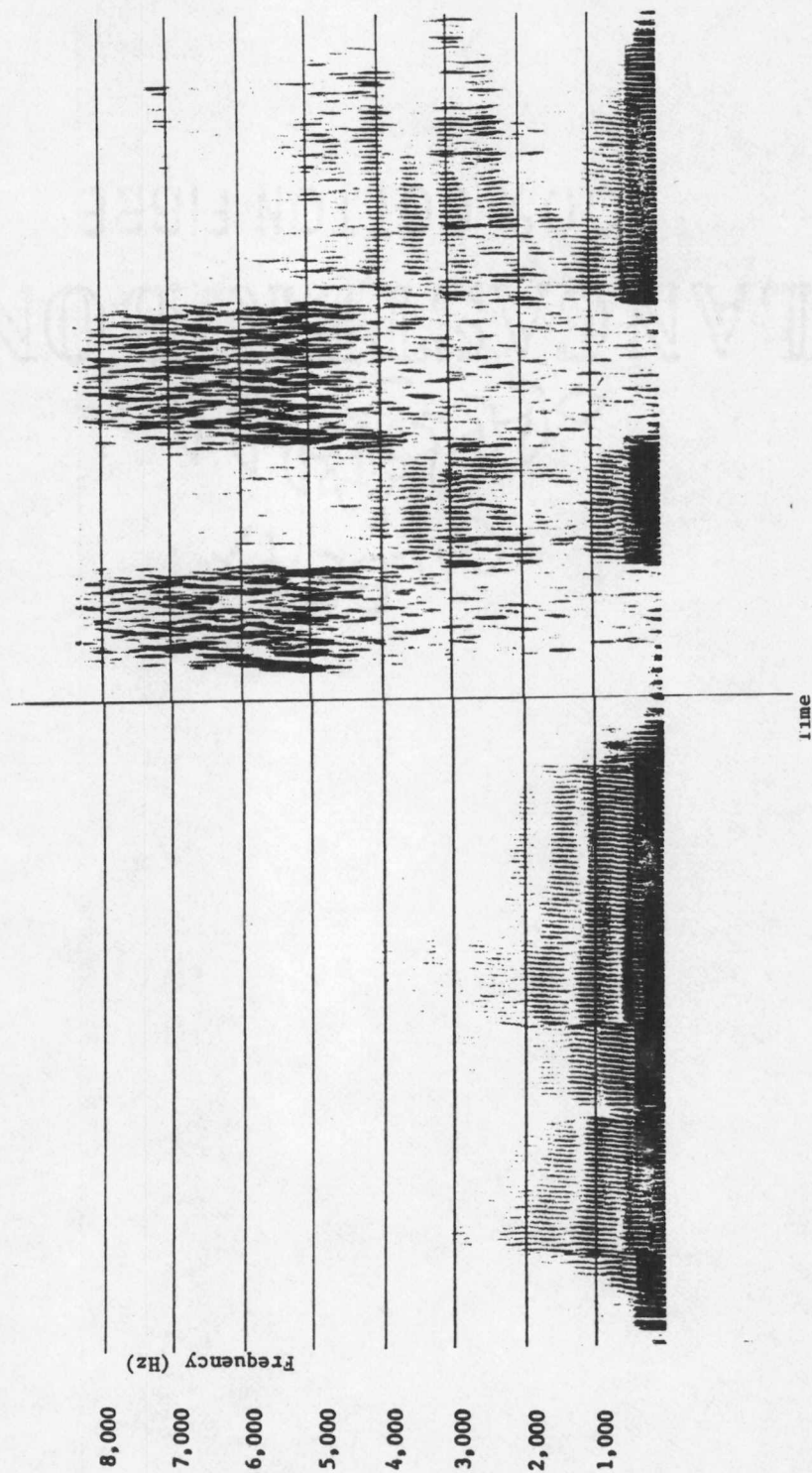
Narrow-band analysis of the homogeneous nonsense syllables /sisi/ - /vovo/ expanded to 2,000 Hz.



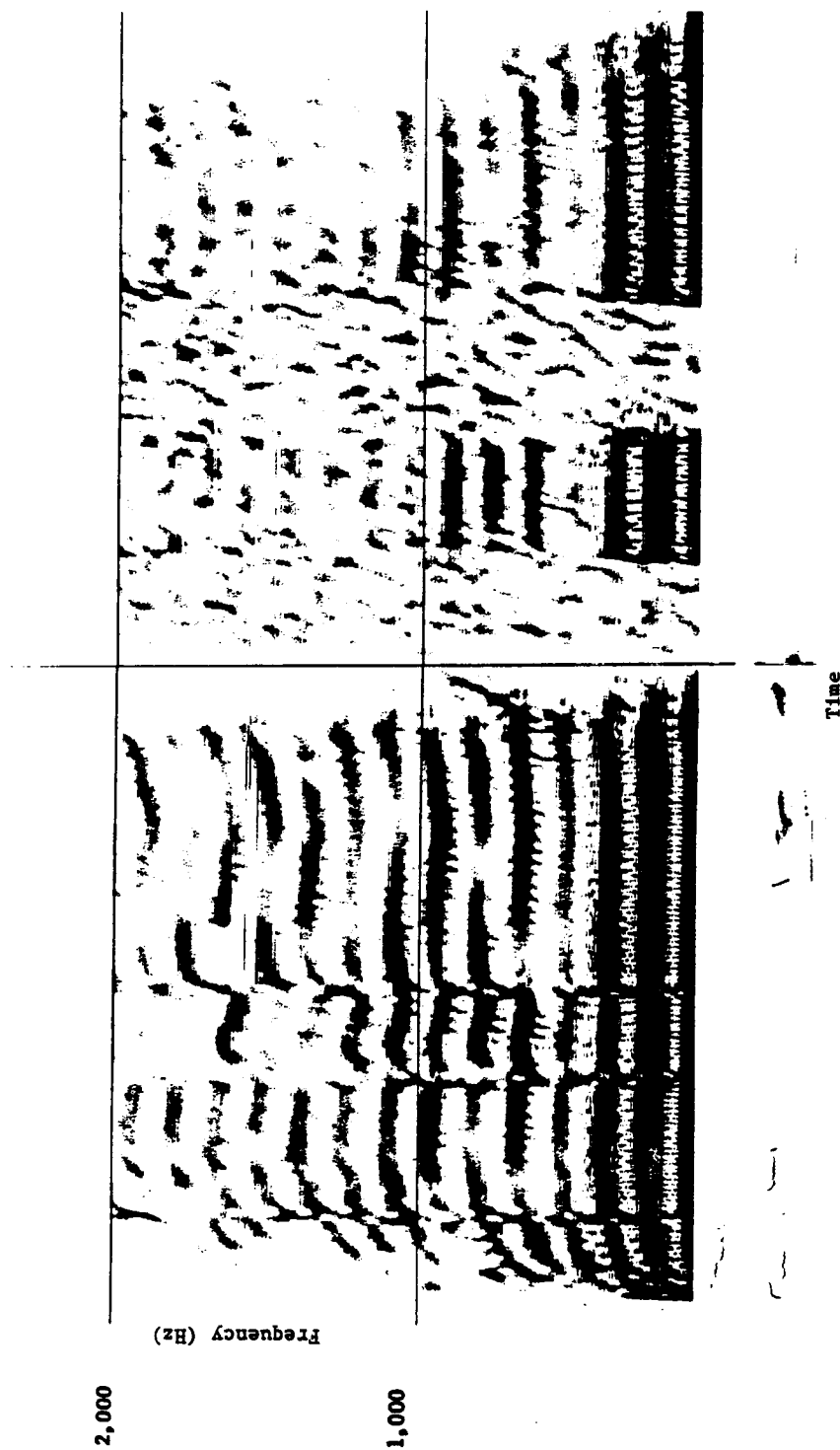
Wide-band analysis of the homogeneous nonsense syllables /mumu/ - /keke/ to 8,000 Hz.



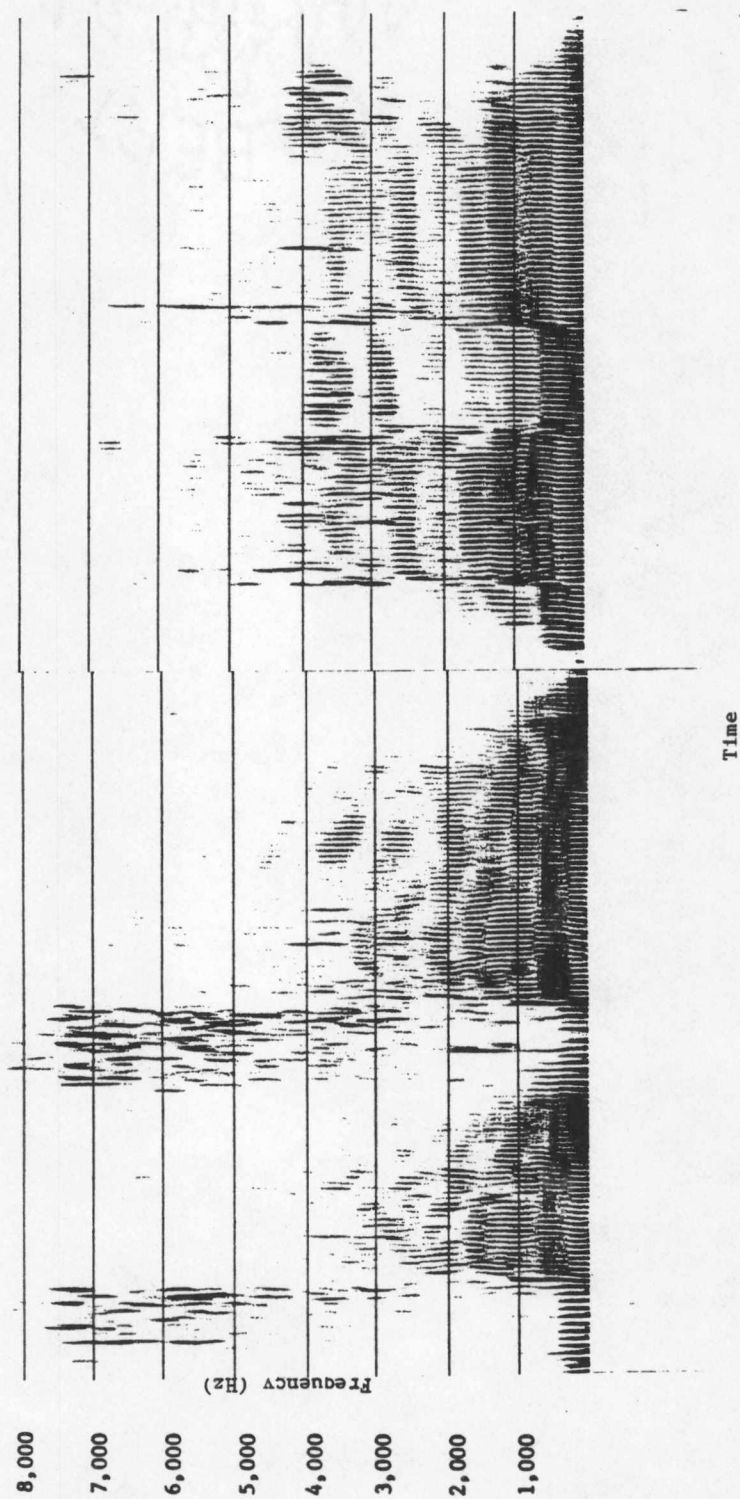
Narrow-band analysis of the homogeneous nonsense syllables /mumu/ - /keke/ expanded to 2,000 Hz.



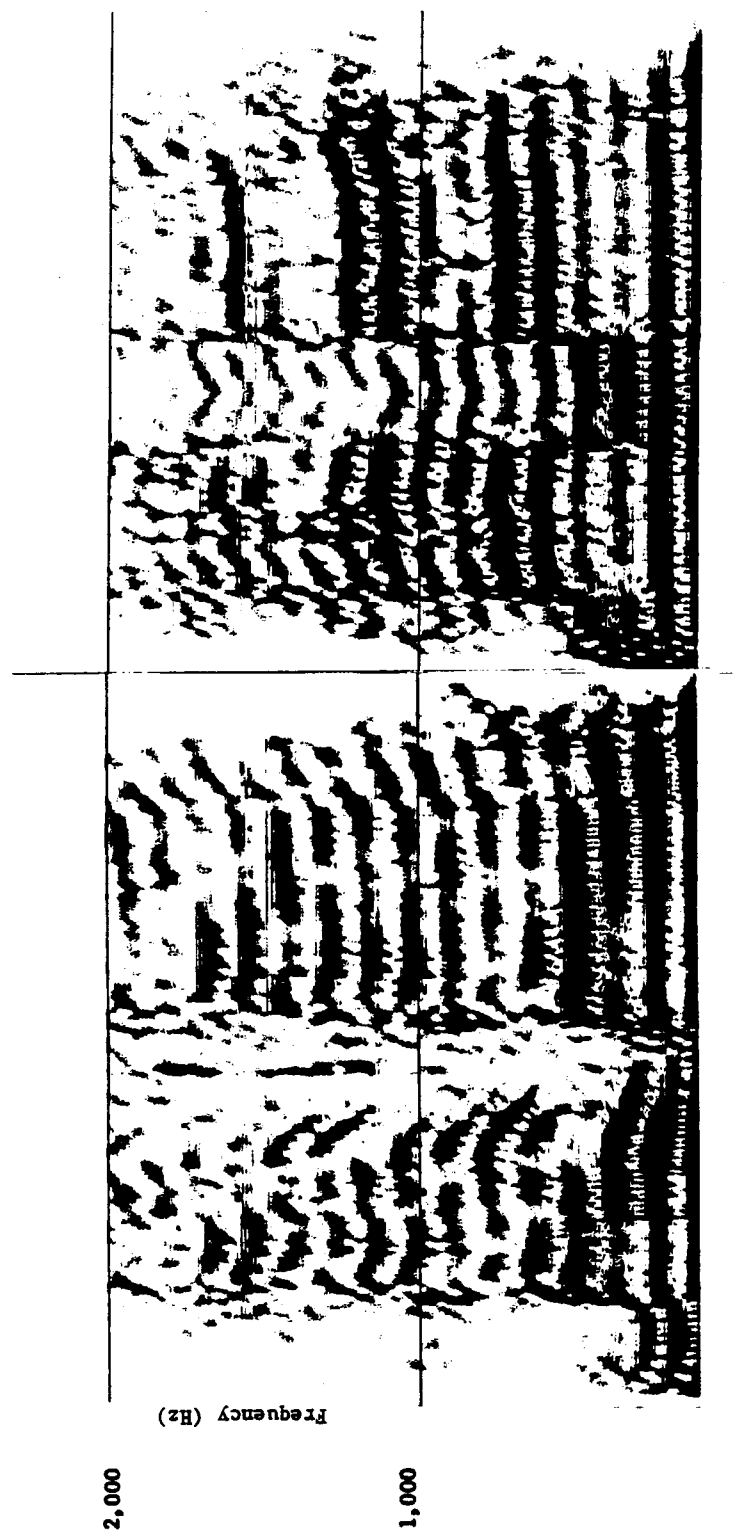
Wide-band analysis for the homogeneous nonsense syllables /mumu/ - /sisi/ to 8,000 Hz.



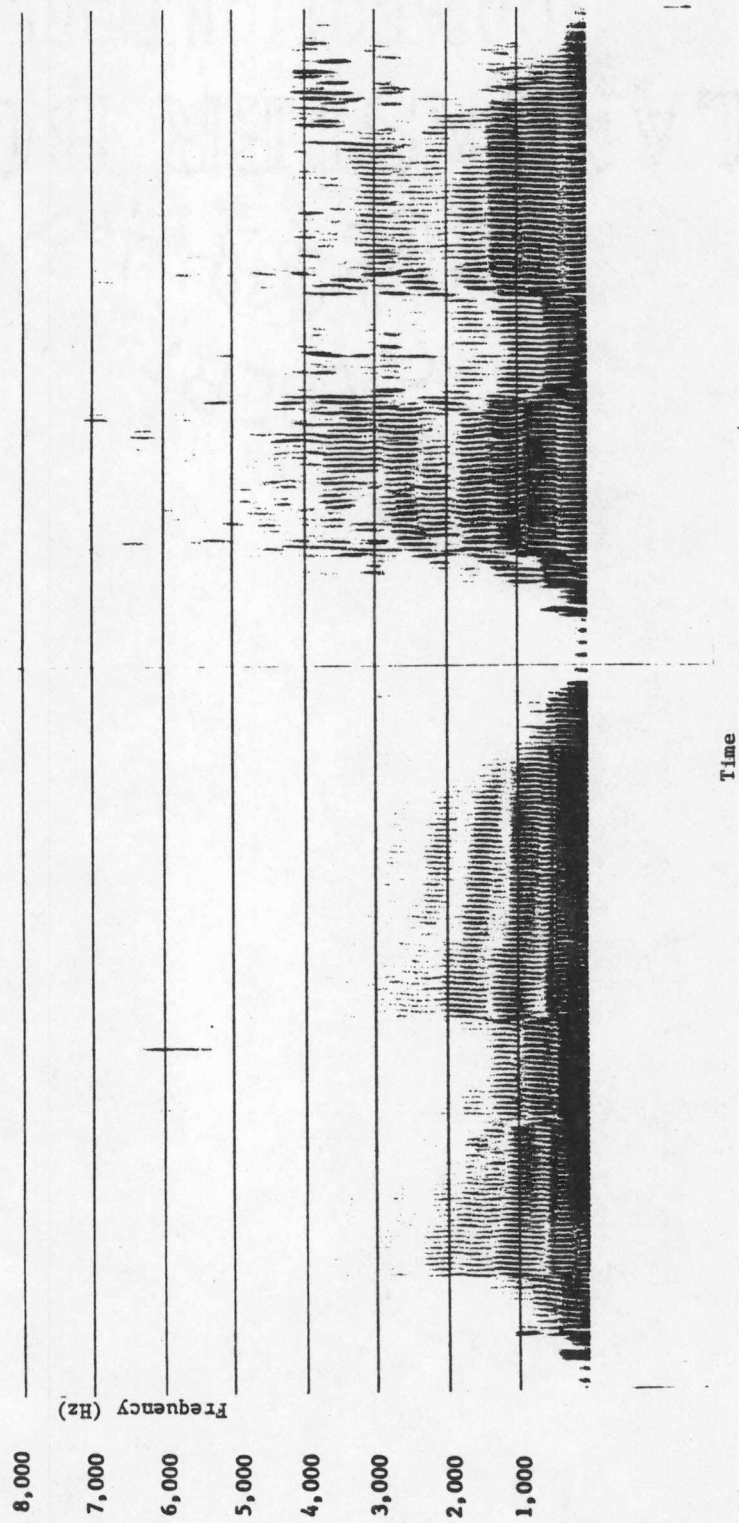
Narrow-band analysis of the homogeneous nonsense syllables /mumu/ - /sisi/ expanded to 2,000 Hz.



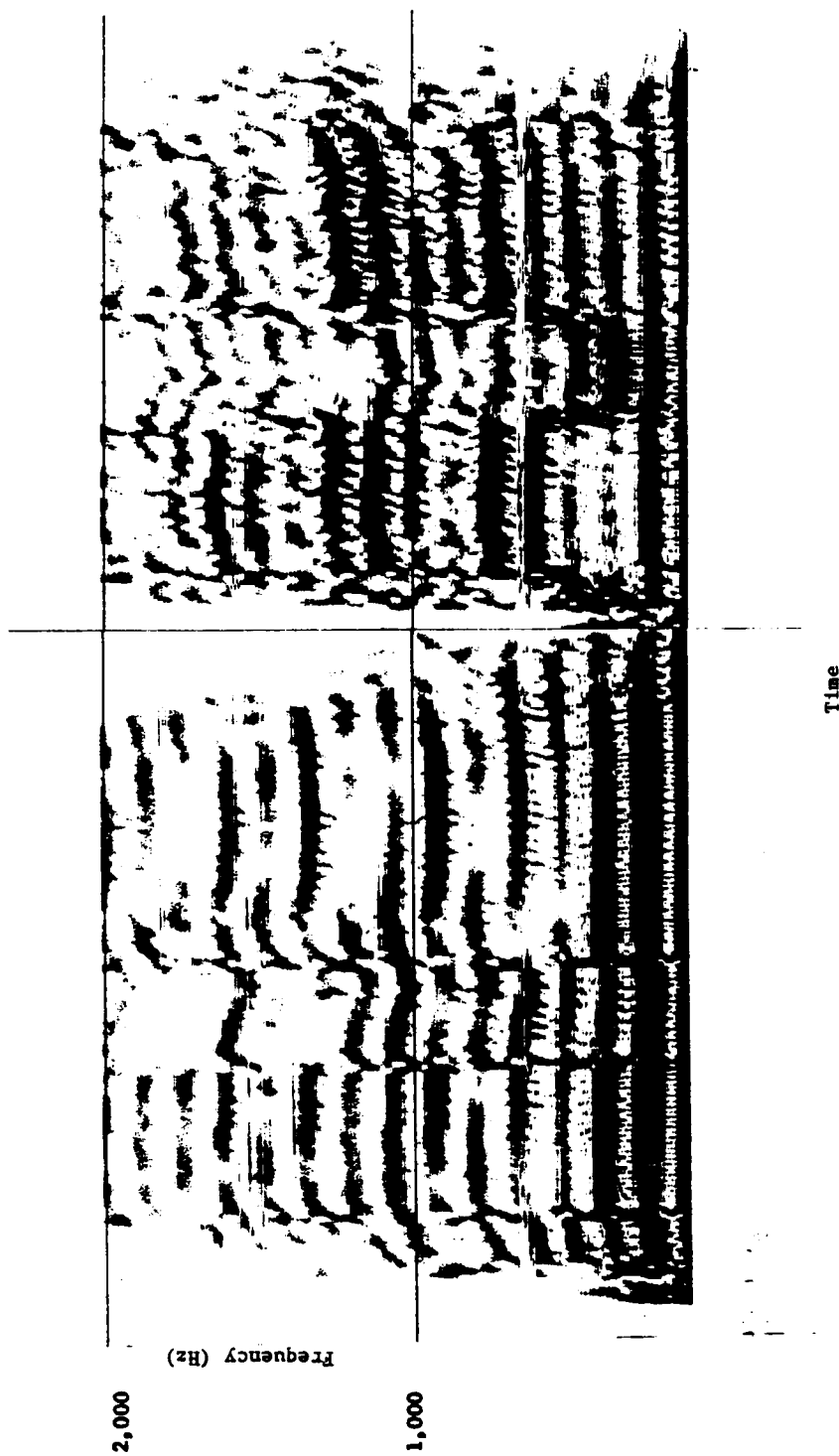
Wide-band analysis of the homogeneous nonsense syllables /vovo/ - /lala/ to 8,000 Hz.



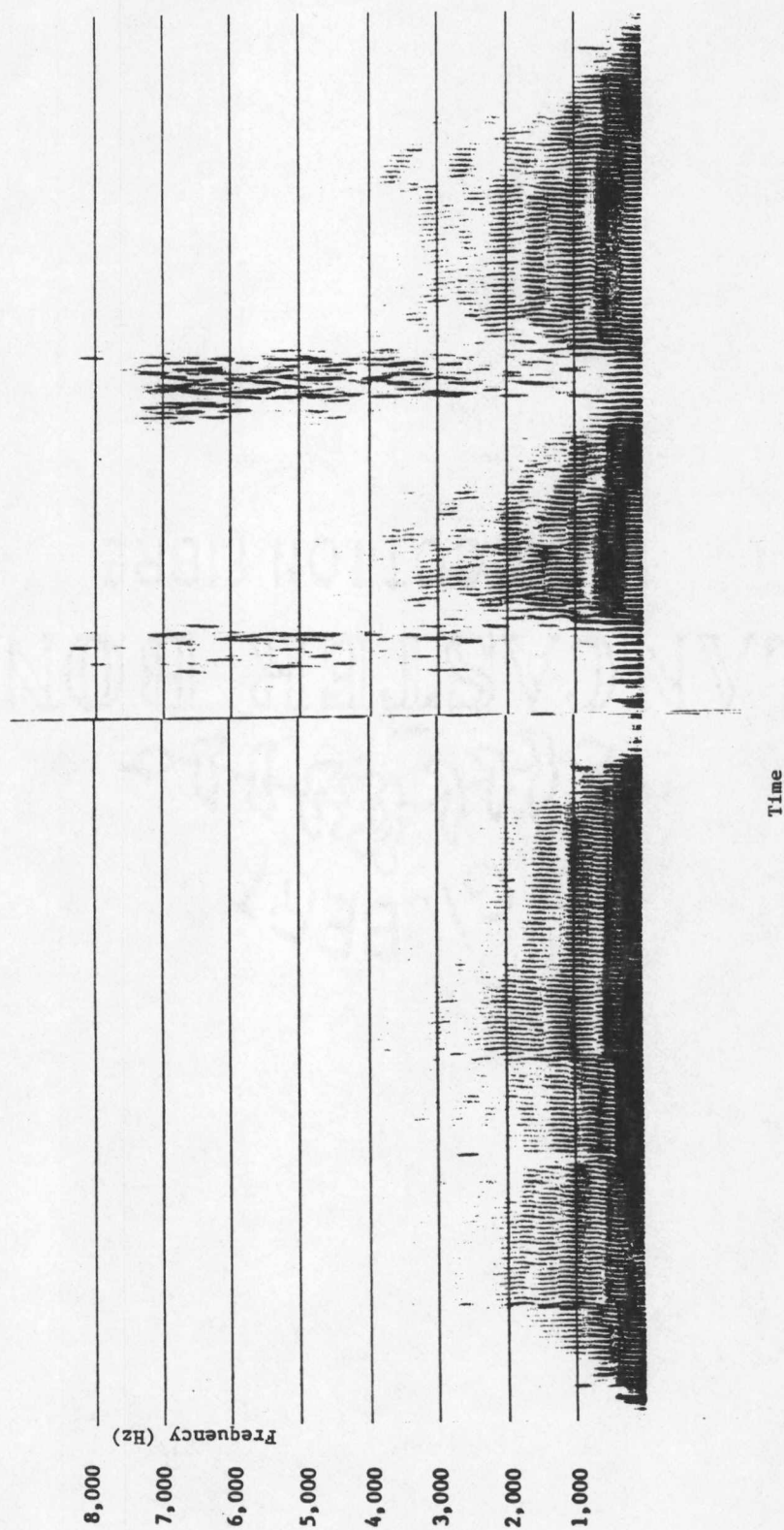
Narrow-band analysis of the homogeneous nonsense syllables /vovo/ - /lqlq/ expanded to 2,000 Hz.



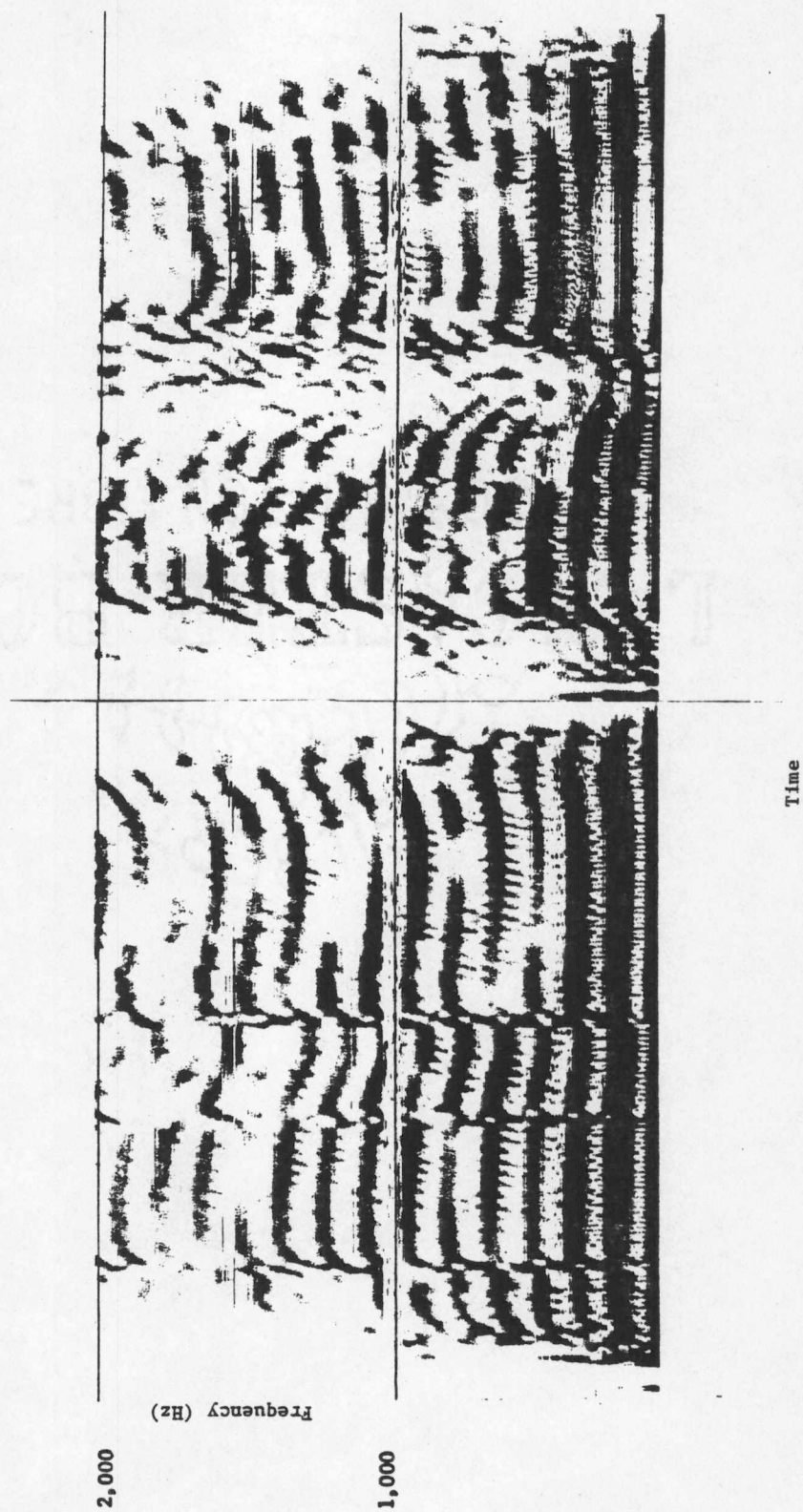
Wide-band analysis of the homogeneous nonsense syllables /mumu/ - /la la/ to 8,000 Hz.



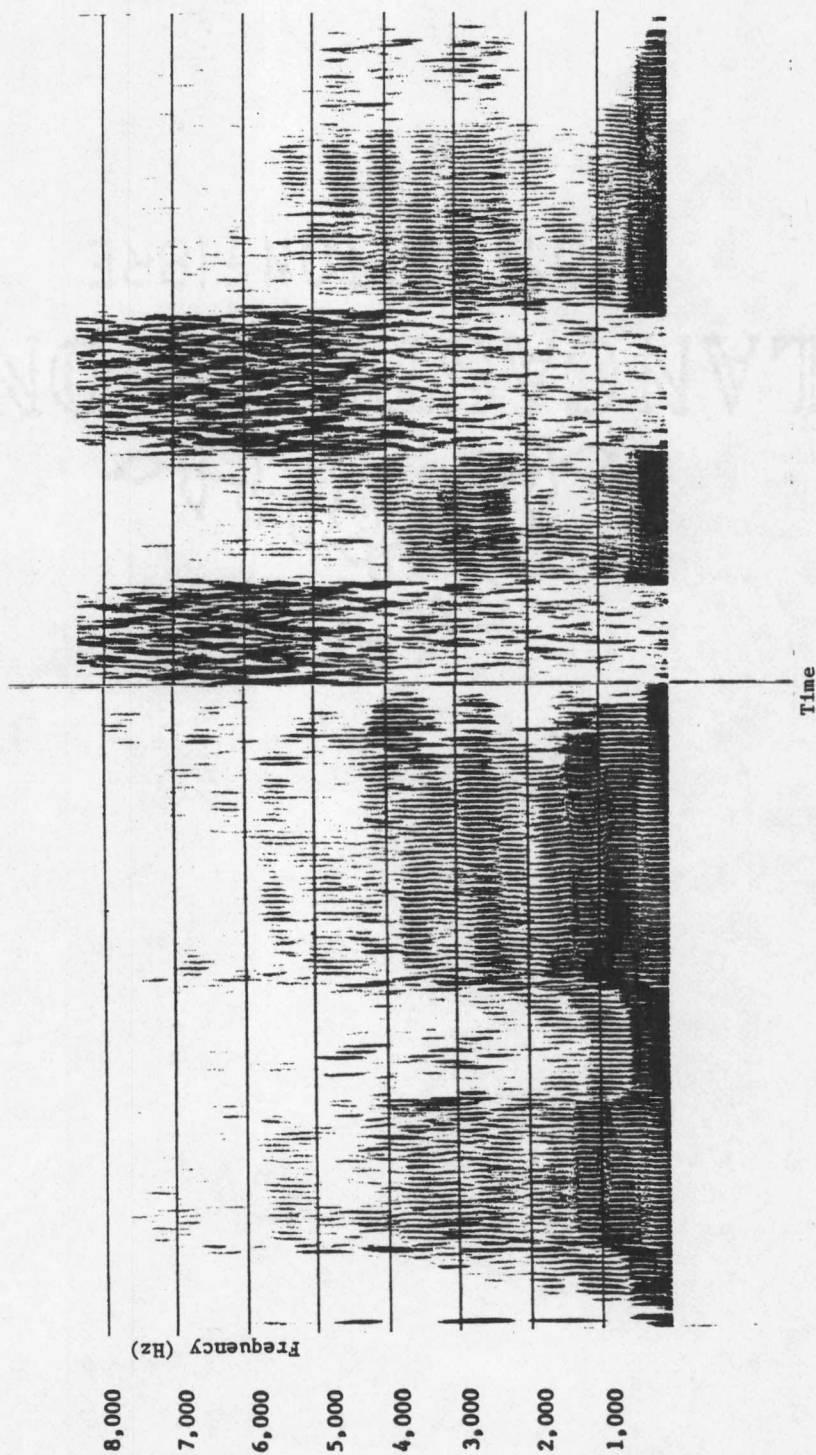
Narrow-band analysis of the homogeneous nonsense syllables /mumu/ - /lɔlɔ/ expanded to 2,000 Hz.



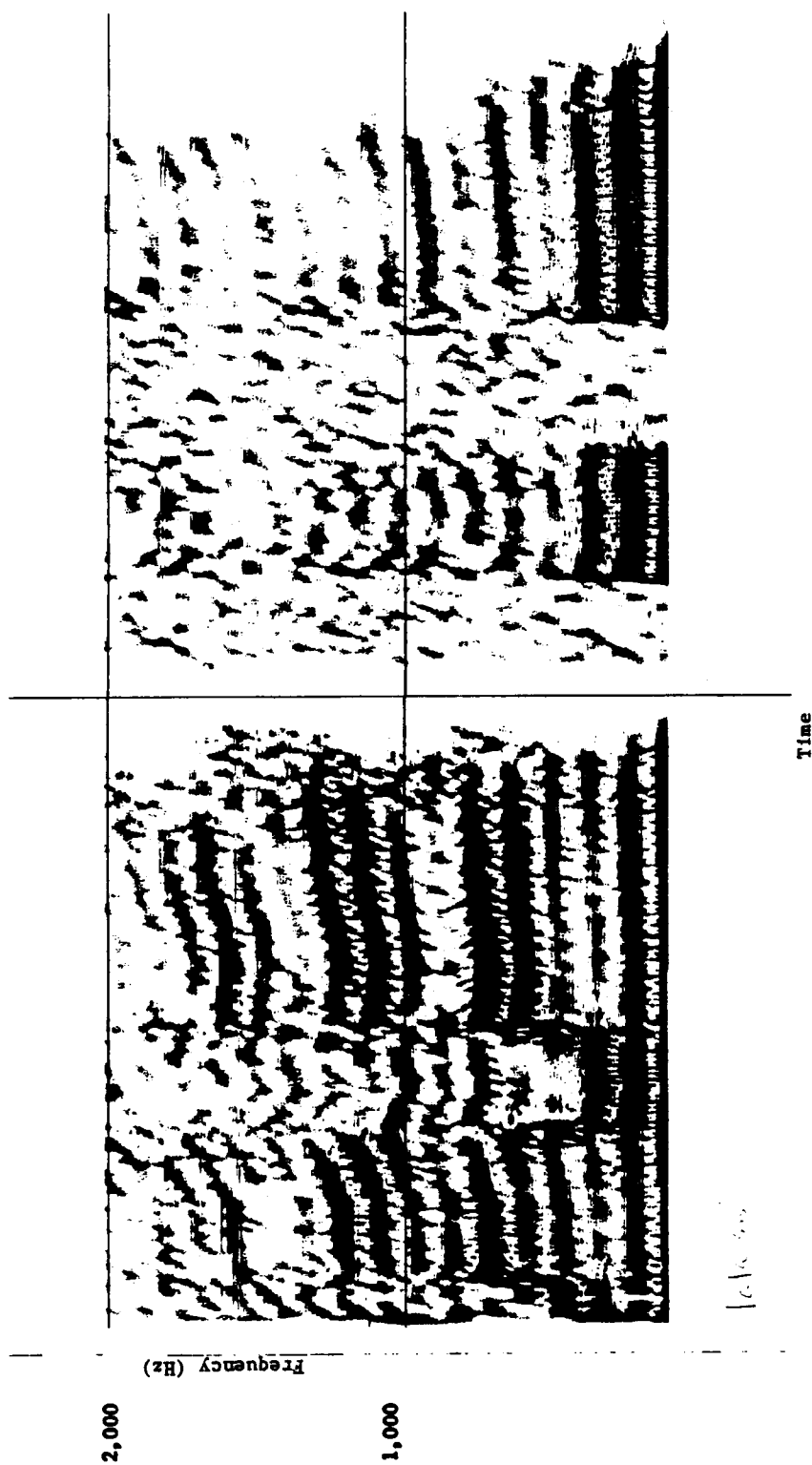
Wide-band analysis of the homogeneous nonsense syllables /mumu/ - /vovo/ to 8,000 Hz.



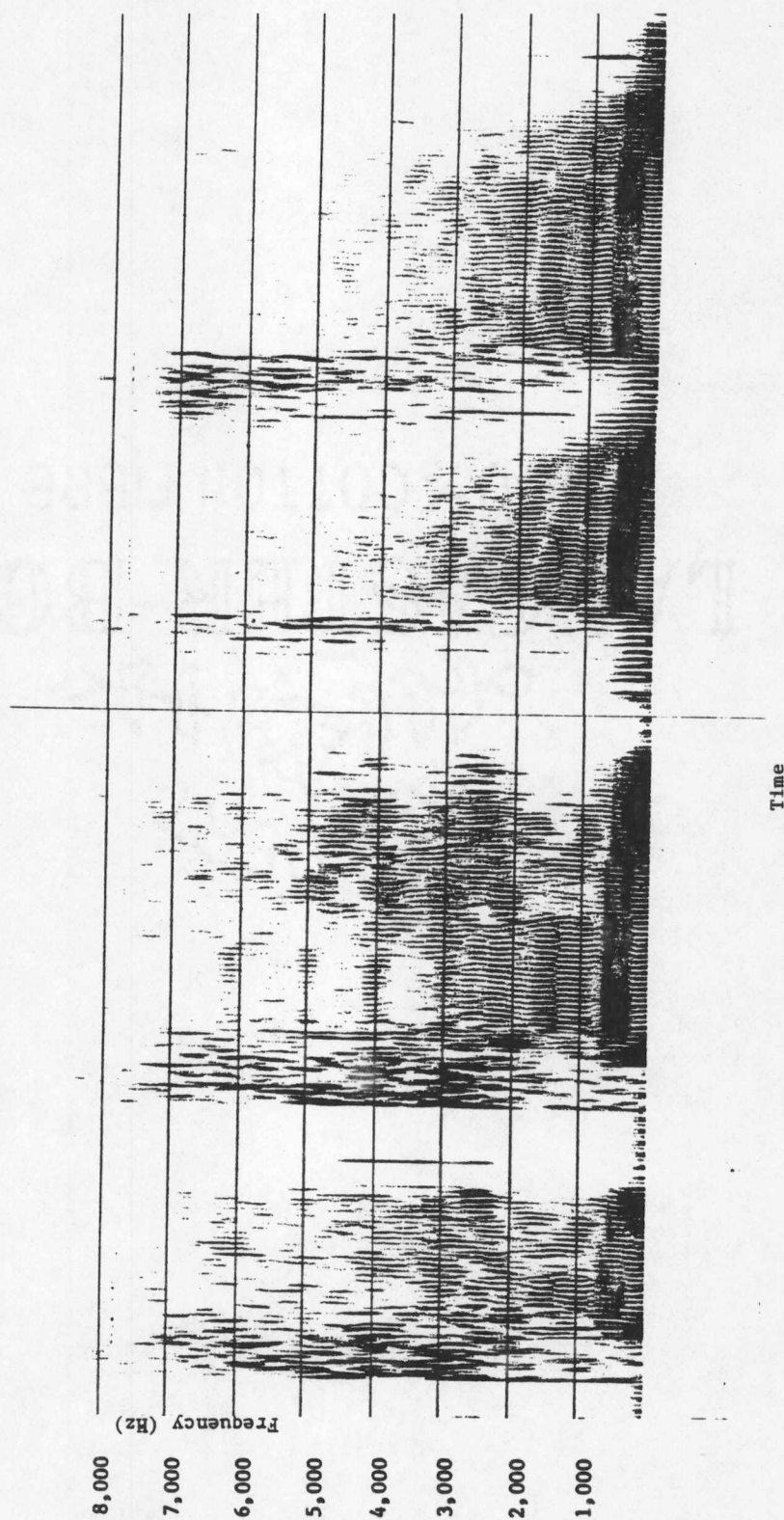
Narrow-band analysis of the homogeneous nonsense syllables /mumu/ - /vovo/ expanded to 2,000 Hz.



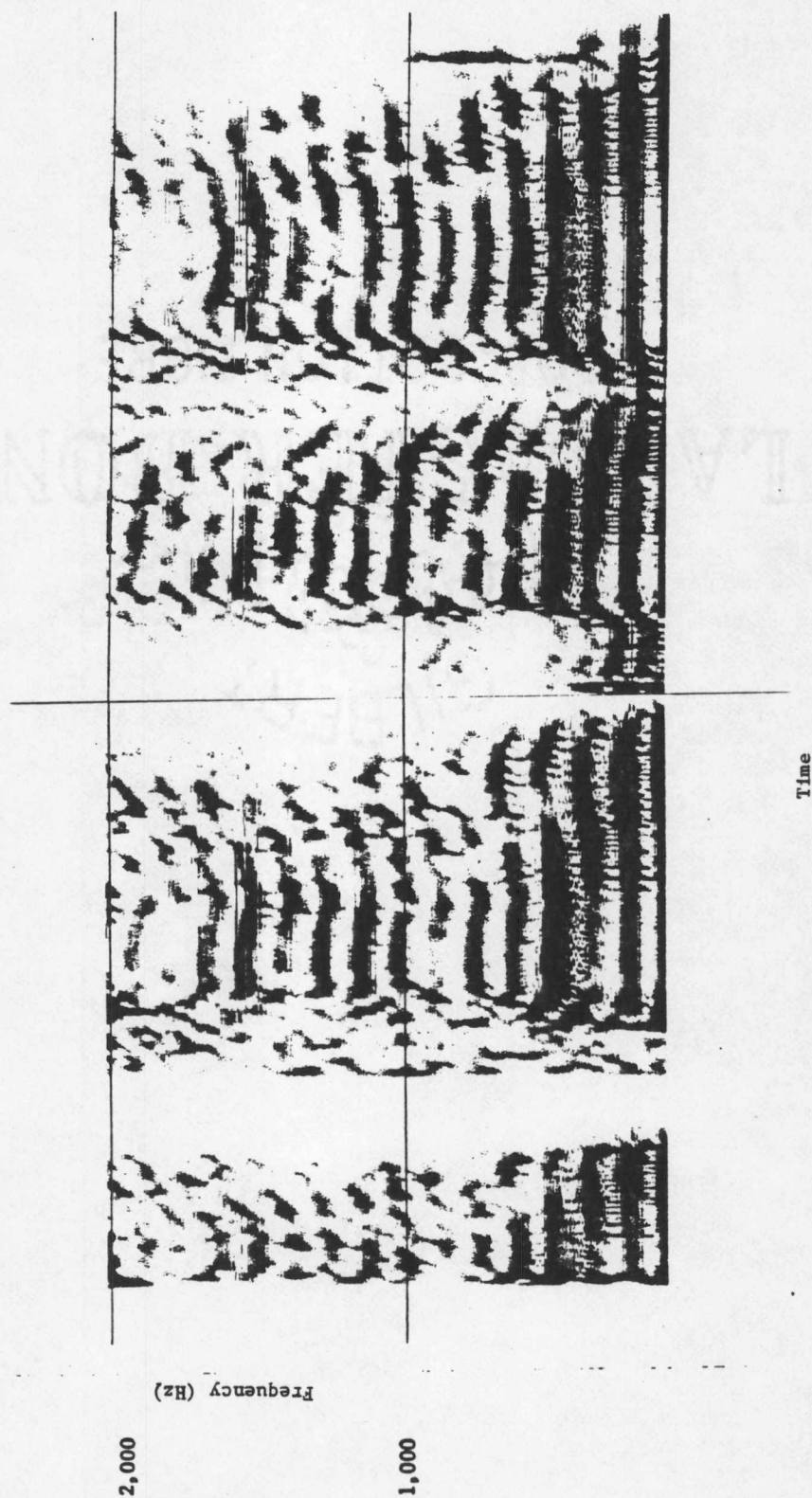
Wide-band analysis of the homogeneous nonsense syllables /lɔlɔ/ - /sisi/ to 8,000 Hz.



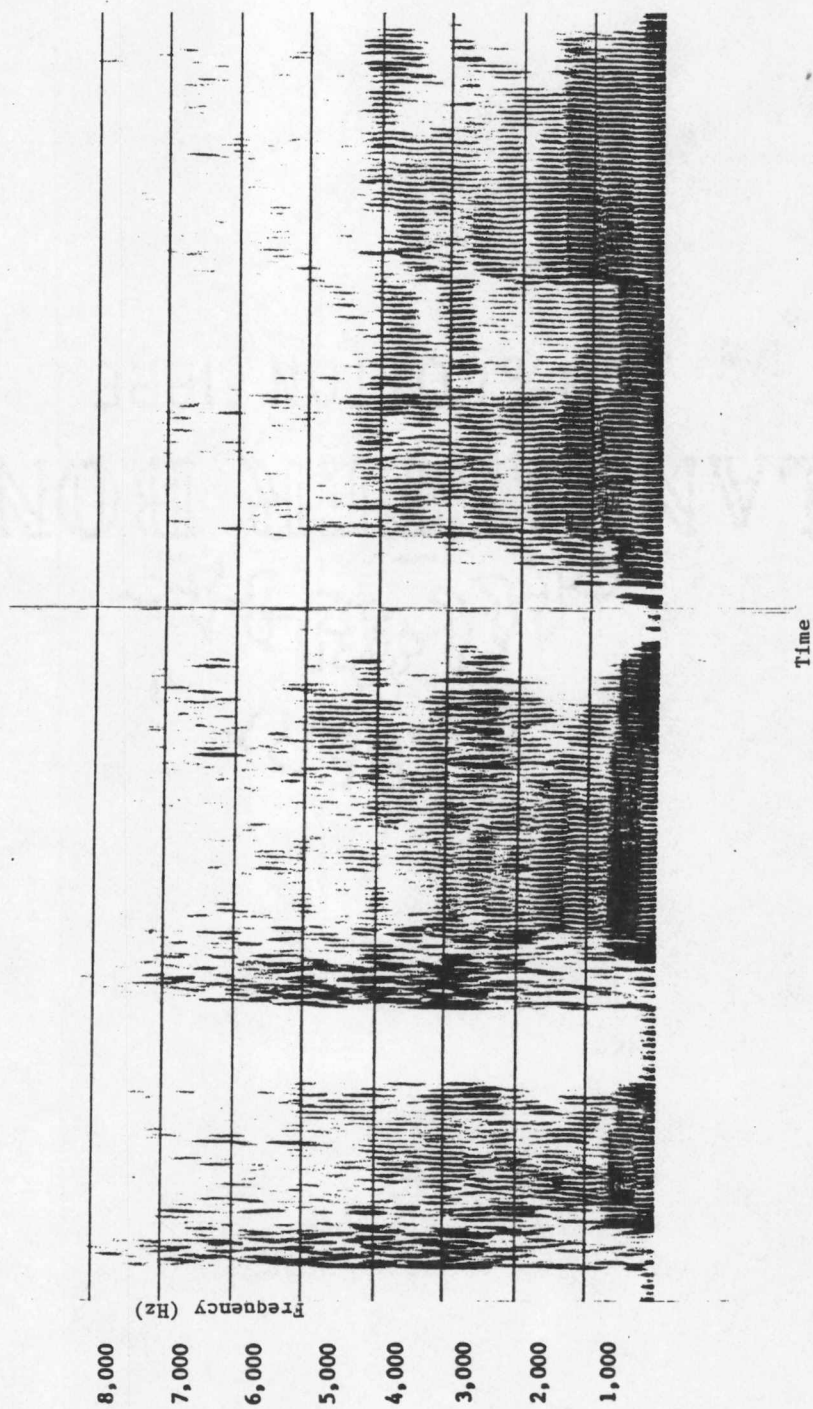
Narrow-band analysis of the homogeneous nonsense syllables /tə'lə/ - /sisi/ expanded to 2,000 Hz.



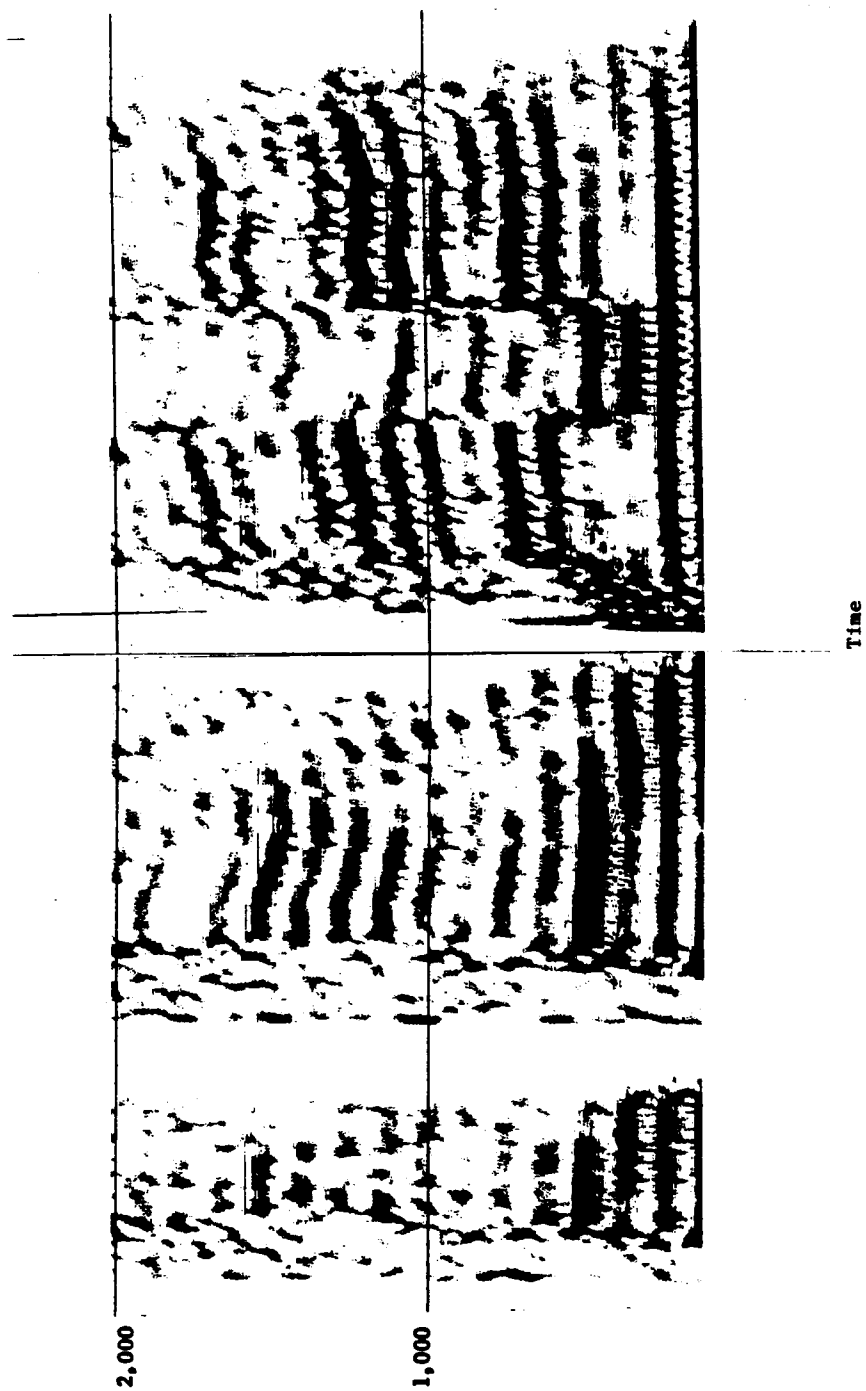
Wide-band analysis of the homogeneous nonsense syllables /keke/ - /vovo/ to 8,000 Hz.



Narrow-band analysis of the homogeneous nonsense syllables /keke/ - /vovo/ expanded to 2,000 Hz.



Wide-band analysis of the homogeneous nonsense syllables /keke/ - /lələ/ to 8,000 Hz.



Narrow-band analysis of the homogeneous nonsense syllables /keke/ - /lələ/ expanded to 2,000 Hz.

APPENDIX I

MATRICES OF SUBJECTS' RESPONSES

Matrix for Responses to Subject 2

	mu	vo	la	ke	si	ai	ai ²
mu	-	1	0	0	0	1	1
vo	0	-	0	0	0	0	0
la	1	1	-	0	0	2	4
ke	1	1	1	-	0	3	9
si	1	1	1	1	-	4	16
						10	28

Matrix for Responses to Subject 5

	mu	vo	la	ke	si	ai	ai ²
mu	-	0	0	0	0	0	0
vo	1	-	0	0	0	1	1
la	1	1	-	0	0	2	4
ke	1	1	1	-	0	3	9
si	1	1	1	1	-	4	16
						10	30

Matrix for Responses to Subject 6

	mu	vo	la	ke	si	ai	ai ²
mu	-	0	0	0	0	0	0
vo	1	-	0	0	0	1	1
la	1	1	-	0	0	2	4
ke	1	1	1	-	0	3	9
si	1	1	1	1	-	4	16
						10	30

Matrix for Responses to Subject 7

	mu	vo	la	ke	si	ai	ai ²
mu	-	0	0	0	0	0	0
vo	1	-	0	0	0	1	1
la	1	1	-	0	0	2	4
ke	1	1	1	-	0	3	9
si	1	1	1	1	-	4	16
						10	30

Matrix for Responses to Subject 9

	mu	vo	la	ke	si	ai	si ²
mu	-	0	0	0	0	0	0
vo	1	-	0	0	0	1	1
la	1	1	-	0	0	2	4
ke	1	1	1	-	0	3	9
si	1	1	1	1	-	4	16
						10	30

Matrix for Responses to Subject 12

	mu	vo	la	ke	si	ai	ai ²
mu	-	0	0	0	0	0	0
vo	1	-	0	0	0	1	1
la	1	1	-	0	0	2	4
ke	1	1	1	-	0	3	9
si	1	1	1	1	-	4	16
						10	30

VITA

The author graduated in June 1977 from Miramonte High School in Orinda, California. She received Bachelor of Arts degrees with Honors in both Psychology and Speech and Hearing Sciences in June 1982 from the University of California at Santa Barbara. As an undergraduate, she received a President's Undergraduate Fellowship for Research to fund her project, "Attitudes Toward Geriatric Hearing Aid Wearers," which she presented at the 1981 ASHA convention in Los Angeles. She entered the MA Program in Audiology at The University of Tennessee, Knoxville in June 1983. She was accepted into the Ph.D. program in Speech and Hearing Sciences (UTK) in October 1984. She served as a Graduate Assistant in the Department of Audiology and Speech Pathology (UTK) from March 1985 to the present. She has published six articles in professional journals and presented 19 papers, poster sessions, and exhibits at state, national, and international meetings.