

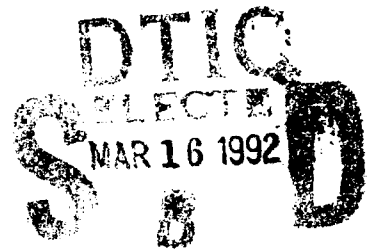


2

NAVAL POSTGRADUATE SCHOOL Monterey, California



THESIS



**Determination of Optimal Training
Methodologies for Discrete/Dependent Speech
Recognition (SR) Systems**

by

Mark C. Rhoads

March 1992

Thesis Advisor:

Gary K. Poock

Approved for public release; distribution is unlimited.

92-06538



**Best
Available
Copy**

Unclassified

SECURITY CLASSIFICATION OF THIS PAGE

REPORT DOCUMENTATION PAGE

1a REPORT SECURITY CLASSIFICATION Unclassified		1b. RESTRICTIVE MARKINGS	
2a SECURITY CLASSIFICATION AUTHORITY		3. DISTRIBUTION AVAILABILITY OF REPORT Approved for public release; distribution is unlimited.	
2b DECLASSIFICATION DOWNGRADING SCHEDULE			
4 PERFORMING ORGANIZATION REPORT NUMBER(S)		5. MONITORING ORGANIZATION REPORT NUMBER(S)	
6a NAME OF PERFORMING ORGANIZATION Naval Postgraduate School	6b OFFICE SYMBOL (If Applicable)	7a. NAME OF MONITORING ORGANIZATION Naval Postgraduate School	
6c ADDRESS (city, state, and ZIP code) Monterey, CA 93943-5000		7b. ADDRESS (city, state, and ZIP code) Monterey, CA 93943-5000	
8a NAME OF FUNDING SPONSORING ORGANIZATION	6b OFFICE SYMBOL (If Applicable)	9. PROCUREMENT INSTRUMENT IDENTIFICATION NUMBER	
8c ADDRESS (city, state, and ZIP code)		10. SOURCE OF FUNDING NUMBERS	
		PROGRAM ELEMENT NO	PROJECT NO
		TASK NO.	WORK UNIT ACCESSION NO.
11 TITLE (Include Security Classification) Determination of Optimal Training Methodologies for Discrete Dependent Speech Recognition (SR) Systems (Unclassified)			
12 PERSONAL AUTHOR(S) LCDR Mark C Rhoads			
13a TYPE OF REPORT Master's Thesis	13b TIME COVERED FROM OCT 90 TO MAR 92	14. DATE OF REPORT (year, month, day) March 1992	15. PAGE COUNT 30
16. SUPPLEMENTARY NOTATION The views expressed in this thesis are those of the author and do not reflect the official policy or position of the Department of Defense or the U.S. Government			
17. COSATI CODES		18. SUBJECT TERMS (continue on reverse if necessary and identify by block number)	
FIELD	GROU	SUBGROUP	
		Speech recognition; training methodologies; experimental results; conclusions	
19. ABSTRACT (Continue on reverse if necessary and identify by block number) A research experiment was conducted to determine whether various combinations of training methodologies and speaking voices would affect recognition accuracies amongst unique speaker dependent Speech Recognition (SR) systems. The experiment used a SR system (VOTAN VTR 6050II) which is based on VOTAN (proprietary) technology. Ten subjects trained five different voice patterns each and conducted four natural voice tests to compile statistics about the recognition accuracy for each pattern. Two patterns (natural voice and declarative voice) were retested using a declarative voice. The experiment was successful and demonstrated that different combinations of training methodologies and speaking voices can significantly affect the performance of unique discrete dependent SR systems. This thesis discusses the research methodology, reviews and analyzes the data collected, and states conclusions drawn about the particular dependent SR system used in the experiment.			
20. Distribution/Availability of Abstract ☒ unclassified/unlimited ☐ same as rpt ☐ DTIC users		21. ABSTRACT SECURITY CLASSIFICATION Unclassified	
22a. NAME OF RESPONSIBLE INDIVIDUAL Gary K. Poock		22b. TELEPHONE (Include Area Code) 408-646-2636	22c. OFFICE SYMBOL Code OR PK

DD FORM 1473, 84 MAR

83 APR edition may be used until exhausted

SECURITY CLASSIFICATION OF THIS PAGE

All other editions are obsolete

Unclassified

Approved for public release; distribution is unlimited.

Determination of Optimal Training Methodologies
for Discrete/Dependent Speech Recognition (SR) Systems
by

Mark C. Rhoads
Lieutenant Commander, United States Navy
B.S., University of Kansas, 1978

Submitted in partial fulfillment of the requirements for
the degree of
MASTER OF SCIENCE IN INFORMATION SYSTEMS
from the

NAVAL POSTGRADUATE SCHOOL
March 1992

Author:

Mark Charles Rhoads
Mark Charles Rhoads

Approved by:

Gary K. Poock
Gary K. Poock, Thesis Advisor

Kenneth W. Thomas
Kenneth W. Thomas, Second Reader

David R. Whipple
David R. Whipple, Chairman, Department of Administrative
Sciences

ABSTRACT

A research experiment was conducted to determine whether various combinations of training methodologies and speaking voices would affect recognition accuracies amongst unique speaker dependent speech recognition (SR) systems. The experiment used a SR system (VOTAN VTR 6050II) which is based on VOTAN (proprietary) technology. Ten subjects trained five different voice patterns each and conducted four natural voice tests to compile statistics about the recognition accuracy for each pattern. Two patterns (natural voice and declarative voice) were retested using a declarative voice.

The experiment was successful and demonstrated that different combinations of training methodologies and speaking voices can significantly affect the performance of unique discrete dependent SR systems. This thesis discusses the research methodology, reviews and analyzes the data collected, and states conclusions drawn about the particular dependent SR system used in the experiment.

Accession For	
NTIS GRA&I	☒
DTIC TAB	☐
Unannounced	☐
Justification	
By	
Distribution/	
Availability Codes	
Dist	Avail and/or Special
A-1	

TABLE OF CONTENTS

I. INTRODUCTION	1
A. BACKGROUND	1
B. PROBLEM	2
C. SCOPE OF THE THESIS	2
D. LIMITATIONS	3
II. EXPERIMENT PROCEDURE	4
A. SUBJECTS	4
B. SR SYSTEM	4
C. EXPERIMENT DESIGN	5
D. PROCEDURE	6
1. Training	6
2. Testing	7
E. INDEPENDENT AND DEPENDENT VARIABLES	8
III. RESULTS	9
A. OVERVIEW	9
1. Analysis of Variance	9
2. Impact of Variables	10
a. 'Subject' Variable	10
b. 'Trial' Variable	10
c. 'Pattern' Variable	12
d. 'Voice' Variable	12
B. DISCUSSION	17
IV. CONCLUSIONS	19
REFERENCES	21
APPENDIX A	22
INITIAL DISTRIBUTION LIST	23

LIST OF TABLES

TABLE I ANALYSIS OF VARIANCE SUMMARY TABLE USING NATURAL VOICE INPUT AGAINST FIVE TYPES OF REFERENCE PATTERNS	10
TABLE II ANALYSIS OF VARIANCE SUMMARY TABLE USING DECLARATIVE VOICE INPUT AGAINST TWO TYPES OF REFERENCE PATTERNS	11
TABLE III DUNCAN'S MULTIPLE RANGE TEST FOR VARIABLE: ACCURACY (NATURAL VOICE INPUT)	16
TABLE IV DUNCAN'S MULTIPLE RANGE TEST FOR VARIABLE: ACCURACY (DECLARATIVE AND NATURAL PATTERNS)	17

LIST OF FIGURES

Figure 1	Subject vs Pattern Average Accuracy, Subjects 1-5	13
Figure 2	Subject vs Pattern Average Accuracy, Subjects 6-10	14
Figure 3	Pattern vs Voice Average Accuracy	15
Figure 4	Effect of Voice on Average Performance	18

I. INTRODUCTION

A research experiment was conducted to determine whether various combinations of training methodologies and speaking voices would affect recognition accuracies amongst unique speaker dependent SR systems. The experiment used a SR system (VOTAN VTR 6050II) which is based on VOTAN (proprietary) technology. Ten subjects trained five different voice patterns each and conducted four natural voice tests to compile statistics about the recognition accuracy for each pattern. Two patterns (natural voice and declarative voice) were retested using a declarative voice. Statistics were compiled on the interaction of these independent variables. This thesis discusses the research methodology, reviews and analyzes the data collected, and states conclusions drawn about the particular dependent SR system used in the experiment.

A. BACKGROUND

This experiment was conducted as follow-on research based on a thesis completed in March 1991 by CDR Richard L. Miller. Each SR system's performance is dependent on whether its algorithms can accurately capture an individual's speech characteristics and later match them to spoken words. The Miller thesis sought to determine whether a dependent SR system's word recognition accuracy would vary significantly with the training method used. Miller's research found a definite relationship between training method and recognition accuracy (Miller, 1991).

A common mistake when using SR equipment is talking too meekly to the system. The system can't recognize what it can't hear (Poock, 1990). Failure to

speak loudly enough causes problems not only during system operation but especially during template training. Declarative speech normally eliminates this problem by naturally causing the speaker to raise his voice. The original research was duplicated with the addition of two new voice patterns. Five types of voice patterns were tested using a natural voice input. In addition, the two patterns which performed best in terms of recognition accuracy were retested using a declarative voice input.

B. PROBLEM

Do optimal training methods exist and if so do they differ amongst unique discrete dependent SR systems? Each dependent SR system is individualistic as defined by the type of algorithms it uses to produce voice templates. An optimal training method for one system may not be the best for other systems. Is it possible to quickly determine an optimal training method for each SR system? Natural voice training is an intuitive method to start with but is it optimal or at least "good enough" when compared to other training methods?

If training methods affect recognition accuracy, a logical follow-on question would be: Can how an individual "speaks" to the computer affect a system's performance? Vendors generally recommend training their SR systems in a natural voice but don't discuss how to speak to the computer during operational use. This thesis addresses these questions as they apply to one specific discrete/dependent SR system.

C. SCOPE OF THE THESIS

The objective of the thesis is to determine whether there is any statistically significant difference in performance between five different training

methodologies, while using two speech types to test a specific, dependent SR system. Training methodologies that are the same as those tested during the Miller research will be compared to determine if a common optimal training method exists.

D. LIMITATIONS

Time limitations precluded conducting the experiment on more than one type of dependent SR system. The results herein are system specific and cannot be generalized for *all* dependent SR systems.

II. EXPERIMENT PROCEDURE

A. SUBJECTS

Ten subjects (two female, eight male) participated in this study. One of the female subjects was a civilian. The remaining subjects were military officers who were enrolled at the Naval Postgraduate School in Monterey, California. Some subjects had educational knowledge of SR systems, but none had actual experience using a SR system before this experiment.

B. SR SYSTEM

The SR system chosen was a stand-alone, off-the-shelf product called 'VOTAN VTR 6050II', which is based on VOTAN SR technology. The algorithm used in the VTR 6050II speech drivers is proprietary. The SR system allows manipulation of two parameters: input gain, and acceptance level. The *acceptance level* can be set on a scale of 0-255 and allows comparison of the spoken utterance with a given template to determine if the accuracy of match is equal to or exceeds the chosen level. A level of zero would require a perfect match while a level of 255 would result in any utterance being recognized. The level was set at the vendor's recommendation, of 50 for this experiment (e.g. if the SR system's algorithm determined a value of 50 or less for a utterance match, it would display the word). The *input gain* allows the user to decrease input gain when using the system in a noisy environment. The gain could be adjusted in a range of values 1-5. The noisier the environment the lower the input gain should be. Input gain was set at a value of 2 even though the experiment was

conducted in a sound proof booth. The system displayed warning messages if the input gain was too high or low.

A noise-cancelling, "boom" microphone mounted on a headset was used for voice input to the system.

C. EXPERIMENT DESIGN

Each subject was given instructions on how to train the SR system. A dumb computer monitor displayed the word being trained and warning messages if the input gain was too low high. The VOTAN VTR 6050II voice card has limited memory capacity and can accept up to 50 words at a time if three training passes are made to create each template. The vendor recommended a set of no more than 20 words in order to enhance recognition and response time. The same vocabulary list of 90 words (Appendix A) used in the Miller study was used to create each template. Due to the memory limitations of the voice card, this list was broken into three separate 30 word lists. Each subject conducted three training passes per template to create five voice templates of each word. Pattern #1--'natural'; Pattern #2--'artificial inflection'; and Pattern #3--'rapid-speak'; Pattern #4--'interrogative'; Pattern #5--'declarative'(see the **Testing** section which follows).

Each subject conducted, on four separate occasions, a series of test runs against their templates using a natural voice. One test run against each template was conducted during each trial session (total of five test runs for each trial; 4 trials x 5 templates = 20 test runs for each subject; total of 20 x 10 subjects = 200 trials). Each template was loaded into the SR system in random order and the subjects were instructed to say each word on the vocabulary list one time. The order of the vocabulary words was modified for each trial to create as much

randomness as possible. The subjects were not allowed to view the computer monitor during trial runs and were not aware of which voice template they were speaking against.

Pattern #1 and Pattern #5 were retested using the same format but with both Voice #1--'natural' and Voice #2--'declarative' speech inputs (total of two test runs for each trial; 4 trials x 2 voice inputs x 2 templates = 16 test runs for each subject; total of 16 x 10 subjects = 160 trials).

D. PROCEDURE

1. Training

Acoustic energy which is produced during speech is affected by changes in loudness, pitch, rate of speech, stress and vocal quality (Tiffany, Carrell, 1977). Each of the five types of templates attempt to take advantage of one or more of these speech qualities. A SR system is dependent on distinctive changes in voice characteristics to produce reliable matching of templates to speech inputs. Templates are more reliable if distinctive vocal features can be incorporated to produce them (Dixon, Martin, 1979). The training templates consisted of 90 vocabulary words, repeated three times by each subject (90x3x10 subjects = 2700 utterances). Each subject created their own, unique templates. Pattern #1, #2 and #3 templates were created in the same manner as they were for the Miller study. Pattern #4 (interrogative) had each subject speak each word as if asking a question. This produced an exaggerated upward or downward inflection on each of the three repetitions. An interrogative type statement will naturally produce either an upward or downward inflection at the end of a word (Tiffany, Carrell, 1977). Pattern #5's templates (declarative) were created in the same manner, each subject speaking the words as if giving the computer a command. A

command type utterance seems to involve an enhancement of all of the speech qualities mentioned above.

During training, the VOTAN system allowed the researcher to accept or reject each utterance by a subject. Acceptance was purely subjective except in the case of input gain being too low/high. The system provided no feedback as to the similarity of utterances. After accepting three repetitions of the utterance, the voice template was saved to computer memory disk. These templates were later input into the system's speech analyzer to test for recognition accuracy. The training procedure took approximately 90 minutes for each subject to train all five voice patterns.

2. Testing

Testing began approximately one week after all subjects had completed creating their templates. Each of the 10 subjects initially conducted four trials each using a natural speaking voice. A trial consisted of five test runs (one for each template). The natural and declarative voice templates were retested using a declarative speaking voice. Testing was made as random as possible. Templates were loaded into the SR system in a random order and each subject read through a corresponding list of vocabulary words. Six lists of vocabulary words were available for each set of 30 words. Words were arranged randomly on each list and each subject was directed to select a different list during each of the four trials. Subjects weren't allowed to know which template was loaded and were not allowed to view the monitor during testing.

During each trial, statistics were recorded as to number of correct recognitions, misrecognitions and nonrecognitions (for the purposes of this thesis.

misrecognitions and nonrecognitions were grouped together and counted as inaccurate recognitions by the SR system).

E. INDEPENDENT AND DEPENDENT VARIABLES

The independent variables were: *pattern* (one, two, three, four and five), *trial* (one through four), *voice* (one, and two) and *subjects* (1-10). The dependent variable was *accuracy*.

III. RESULTS

A. OVERVIEW

This section describes the results of the experiment. The analysis of variance and Duncan Range tests were performed using the arc sin transformation of relative difference scores to stabilize the variance of the error terms (Neter and Wasserman, 1974). The SR recognition accuracy figures that appear in charts, however, are expressed as percentages and are untransformed.

From a statistician's viewpoint, the null hypothesis in this experiment was that all training methods for a dependent SR system would result in equivalent performance.

1. Analysis of Variance

Table I and Table II present respectively the three-way and four-way analysis of variance summary tables for recognition accuracy (arc sin transformation of raw data). F-ratios in Table I indicate that while the 'pattern' and 'subject' variables and their combination had significant effects on the results, 'trials' had no appreciable effect. The F-ratios in Table II again show that 'trials' had no significant effect on the results while 'pattern,' 'subject,' 'voice' and their two-way interactions did. The three-way interaction of 'subject'- 'pattern'- 'voice' was not significant.

2. Impact of Variables

a. 'Subject' Variable

As expected, variability existed between subjects in regard to which patterns and type voice performed better, however their variance is isolated in this design.

b. 'Trial' Variable

The 'trial' variable had no significant affect in either phase of this study. Words were arranged randomly on each vocabulary list and this apparently eliminated any "learning" by the subjects.

TABLE I

ANALYSIS OF VARIANCE SUMMARY TABLE
USING NATURAL VOICE INPUT AGAINST
FIVE TYPES OF REFERENCE PATTERNS

Source	df	SS	MS	F-ratio	Prob
Pattern	4	458.3693	114.5923	27.07	.0001
Trial	3	3.71140	1.237133	0.29	0.8309
Subj	9	1155.6828	128.4092	30.33	.0001
Pattn.Trial	12	30.9971	2.58309	0.61	0.8296
Pattn.Subj	36	547.6957	15.21377	3.59	.0001
Trial.Subj	27	80.3976	2.9777	0.70	0.8530
Error	108	457.1939	4.2333		
Total	199	2734.0478			

TABLE II

ANALYSIS OF VARIANCE SUMMARY TABLE
 USING DECLARATIVE VOICE INPUT AGAINST
 TWO TYPES OF REFERENCE PATTERNS

Source	df	SS	MS	F-ratio	Prob
Pattern	1	3.5701	3.5701	1.99	0.1701
Trial	3	4.0802	1.3601	0.76	0.5281
Subj	9	201.3841	22.3760	12.45	0.0001
Voice	1	20.3776	20.3776	11.34	0.0023
Pattn.Trial	3	8.2027	2.7342	1.52	0.2315
Pattn.Subj	9	50.6256	5.6251	3.13	0.0103
Trial.Subj	27	35.1517	1.3019	0.72	0.7961
Subj.Voice	9	47.8081	5.3120	2.96	0.0140
Pattn.Voice	1	14.4601	14.4601	8.05	0.0085
Voice.Trial	3	3.2162	1.0721	0.60	0.6227
Subj.Pattn. Voice	9	14.3556	1.5951	0.89	0.5485
Subj.Pattn. Trial	27	50.8292	1.8826	1.05	0.4524
Patn. Voice. Trial	3	2.8927	0.9642	0.54	0.6612
Subj.Voice. Trial	27	47.9557	1.7761	0.99	0.5120
Error	27	48.5192	1.7970		
Totals	159	553.4284			

c. 'Pattern' Variable

The 'pattern' variable has a significant effect on performance, as depicted in Figures 1, 2 and 3. Figures 1 and 2 show the differences in pattern performance for each subject. Figure 3 shows the effect that the interaction of pattern and voice had on performance.. To further isolate and analyze the 'pattern' variable, Duncan's Multiple-Range test was conducted. The results of the test are summarized in TABLES III and IV. Note that there is no significant difference in percent accuracy between the natural and declarative patterns (Pattern #1 vs Pattern #5) when tested with a natural speech input (Table III).

d. 'Voice' Variable

The natural (Pattern #1) and declarative (Pattern #5) patterns were retested using a declarative voice. Figure 3 demonstrates that the interaction of input voice type and pattern type did significantly effect percent accuracy. Table IV shows the Duncan Range analysis of means for the two voice types. A declarative voice (Voice #2) takes advantage of all the positive qualities of spoken speech and seems to improve performance when used as a speech input even though there was no appreciable difference between the natural and declarative patterns using a natural input voice (Voice #1).

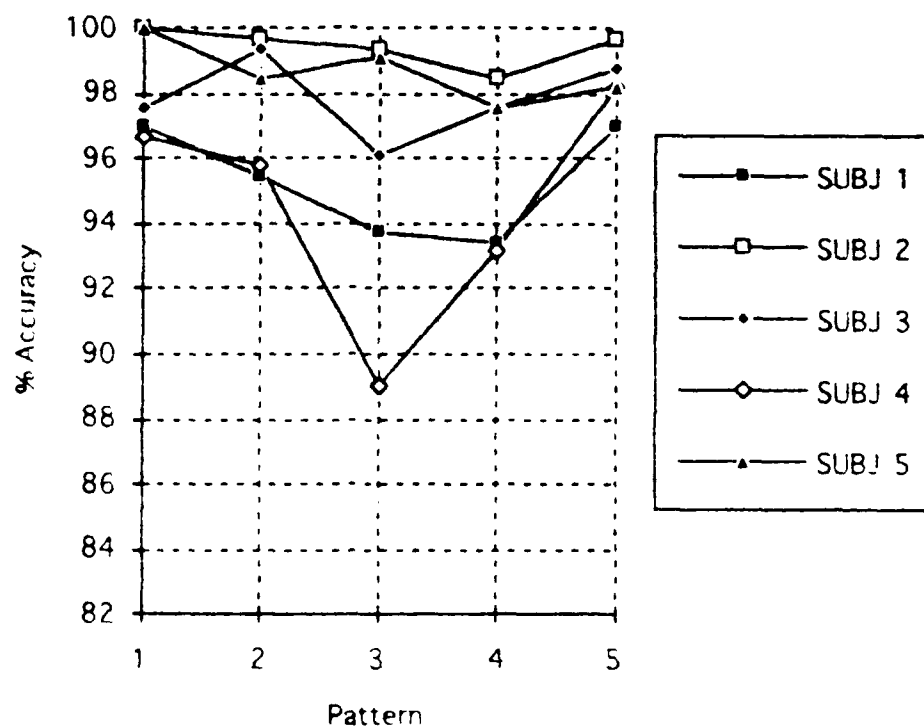


Figure 1. Subject vs Pattern Accuracy, Subjects 1-5
(Patterns: 1 = natural, 2 = artificial inflection, 3 = rapid-speak,
4 = interrogative, 5 = declarative)

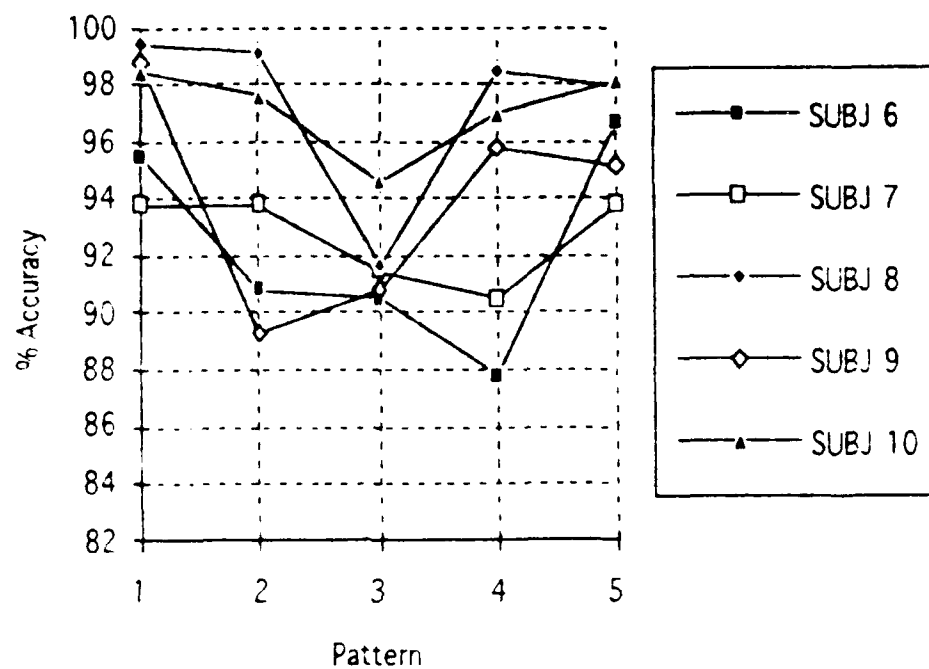


Figure 2. Subject vs Pattern Accuracy, Subjects 6-10
(Patterns: 1 = natural, 2 = artificial inflection, 3 = rapid-speak,
4 = interrogative, 5 = declarative)

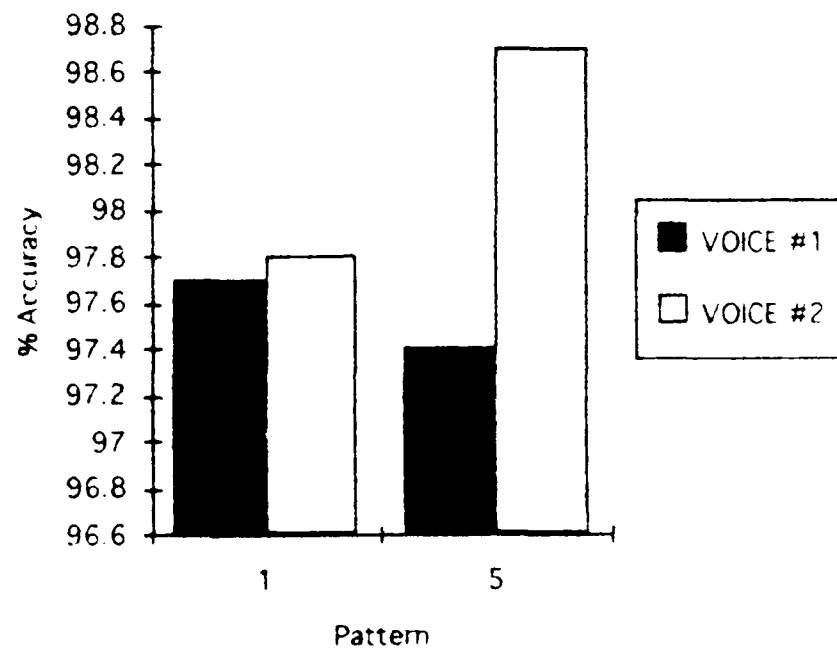


Figure 3. Pattern vs Voice Average Accuracy

(Patterns: 1 = natural, 5 = declarative)

(Voices: 1 = natural, 2 = declarative)

TABLE III

Duncan's Multiple Range Test for Variable : ACCURACY

Natural Voice Input

Alpha= 0.05		df= 108		MSE= 4.2333
Number of Means	2	3	4	5
Critical Range	0.914	0.961	0.991	1.014
Means with the same letter are not significantly different				
Duncan Grouping		Mean	N	PATTERN
	A	97.7275	40	1 (natural)
	A	97.365	40	5 (declarative)
	B	95.94	40	2 (artificial inflection)
	C	94.9925	40	4 (interrogative)
	D	93.63	40	3 (rapid-speak)

TABLE IV

Duncan's Range Test for Variable: ACCURACY
Declarative and Natural Patterns

Alpha= 0.05	df= 27	MSE= 1.7970	
Number of Means	2		
Critical Range	0.4346		
Means with the same letter are not significantly different.			
Duncan Grouping	Mean	N	Voice
A	98.2600	80	2 (declarative)
B	97.5462	80	1 (natural)

B. DISCUSSION

This experiment did evaluate the overall SR accuracy of five training methods by using a natural speaking voice input into the VOTAN VTR 6050II system. Patterns one and five were not significantly different when compared to each other but were appreciably better than the other three patterns (Table III). This supports the Miller study which found that a natural voice pattern performed best. The recommendation in the SR system's documentation was to train the system in a firm, natural voice. The declarative voice pattern was an attempt to interpret these recommendations. The natural and declarative patterns were consistently accurate for all subjects. Patterns two and three did not perform as

well and were not as consistent. The rapid speech pattern in both studies was clearly not as robust as any of the other patterns.

After determining that patterns one and five clearly resulted in more accurate recognitions, the subjects retested patterns one and five using a declarative voice input. As indicated by Figures 3 and 4, the declarative voice input significantly improved the performance both patterns achieved with a natural voice input.

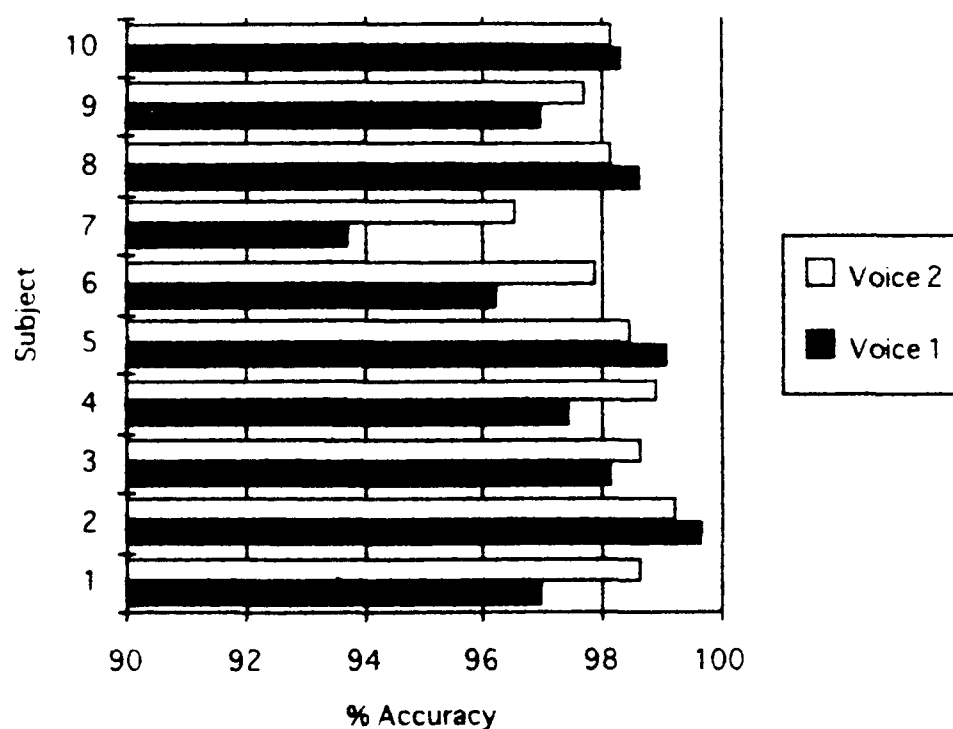


Figure 4. Effect of Voice on Average Performance
(Voices: 1 = natural, 2 = declarative)

IV. CONCLUSIONS

In summary, subjects, as expected impacted performance, but their variance was isolated for this experiment's design. The trial variable had no effect on this study. The effect of pattern, input voice and their interaction did significantly impact performance of the system.

All patterns, with the exception of rapid speech, performed reasonably well. However, the natural and declarative templates clearly achieved the best recognition accuracy. Subjects tended to have difficulty producing the pattern two and four templates. Each subject had several utterances rejected because they weren't able to produce the correct inflection, utterances weren't loud enough, etc. Producing training templates must be an easy, straight-forward and intuitive process if SR systems are to be readily accepted in the market place. Training in a natural voice is an obvious starting point and may produce acceptable results but as demonstrated in both studies, there are a wealth of different methods that could be used. There is not an obvious, or simple way to determine a SR system's optimal training method without conducting experiments similar to this one because each system's algorithms are different.

This experiment demonstrated that recognition accuracy is also dependent on the type of voice used during system operation. Changing from a natural to a declarative voice during testing appreciably improved the system's performance. Declarative utterances are very intuitive to make and generate subtle differences in syllable stress, cadence, inflection and loudness. In this case, a declarative template combined with a natural voice input produced accuracies that were not significantly different from those produced by a natural template and a natural

voice input. However, a declarative template combined with a declarative voice input was significantly better than any pattern or combination that was tested.

Does this mean that all systems should be trained and operated using a declarative voice? Not necessarily because each system is different. Again it's a reasonable method to start with and may produce acceptable or even optimal results depending on the SR system. Manufacturers of SR systems should test their systems using a variety of training methods and input voices to determine the best method for their specific system. They should then give concise and easily understood instructions on the best method to train and use their system. Vague or difficult to grasp directions do little to improve performance of the systems and can actually hinder it. The bottom line is customer satisfaction and a little research and documentation up front can go a long way to improve the acceptance of speech recognition systems.

The Naval Postgraduate School has many different state-of-the-art speech recognition systems and this writer would recommend that support from sponsors be provided to further resolve the questions posed in this thesis. The point of contact at NPS would be this writer's thesis advisor.

REFERENCES

- Dixon, N. Rex and Martin, Thomas B., *Automatic Speech and Speaker Recognition*, IEEE Press, 1979.
- Miller, Richard L., *Training Methodologies for Dependent Speech Recognition (SR) Systems*, Master's Thesis, Naval Postgraduate School, Monterey, California, March, 1991.
- Neter, J. and Wasserman, W., *Applied Linear Statistical Models*, Richard D. Irwin, Inc., 1974.
- Poock, Gary K., Slides from Stockholm, Sweden. International Society of Augmentative and Alternative Communication Conference. Naval Postgraduate School, August, 1990.
- Tiffany, William R. and Carrell, James. *Phonetics Theory and Application*, McGraw-Hill, Inc., 1977.

APPENDIX A

VOCABULARY LIST

ACTIVATE	FIVE	PEAS	TRANSMISSION
ALFA	FOUR	PROBABILITY	TWO
ALTITUDE	FOXTROT	PROCEED	UNIFORM
APPLICATIONS	GALE	PROTOCOL	VICTOR
ASTERISK	GOLD	QUEBEC	VOICE_COMMANDS
ATTACK	GOLF	RAZE	VOICE_HELP
BINGO	HOTEL	RACE	VOICE_OPTIONS
BRAVO	IDENTIFICATION	RECOGNITION	WHISKEY
BUSINESS	INDIA	REFUEL	XRAY
CANCEL	INTERACTIVE	RELOCATE	YANKEE
CHARLIE	JULIET	REPORT	ZERO
CLOSE_WINDOW	KID	ROMEO	ZULU
COMBINATION	KILO	SCRATCH_THAT	
COMMANDER	KIT	SEVEN	
CONTROLLER	LABEL	SIERRA	
COPY	LAUNCH	SIX	
CORPORATION	LIMA	SPEED	
DEACTIVATE	LIST	SOLD	
DELTA	MANEUVER	STATION	
DESIGNATE	MIKE	SUITABILITY	
DETECTION	NINE	SWITCH_APPLICATION	
DISTANCE	NOVEMBER	TALE	
ECHO	ONE	TANGO	
EIGHT	OSCAR	THREE	
ENGINEERING	PAPA	TIME	
EXPRESSWAY	PEACE	TOP_LEVEL	

INITIAL DISTRIBUTION LIST

	No. Copies
1. Library, Code 52 Naval Postgraduate School Monterey, California 93943-5100	2
2. Gary K. Poock, Code ORPK Naval Postgraduate School Monterey, CA 93943	3
3. Kenneth Thomas, Code AS TH Naval Postgraduate School Monterey, CA 93943	1
4. LCDR Mark C. Rhoads 127 Steeplechase Dr. Doylestown, PA 18901	4
5. Defense Technical Information Center Cameron Station Alexandria, VA 22304-6145	2