

AD-A242 983



NAVAL POSTGRADUATE SCHOOL
Monterey, California



DTIC
ELECTE
DEC 6 1991
S C D

THESIS

**Training Methodologies for Dependent
Speech Recognition (SR) Systems**

by

Richard L. Miller

March 1991

Thesis Advisor:

Gary K. Poock

Approved for public release; distribution is unlimited.

91-17169



91 12 048

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. DCLASSIFICATION/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 NAVAIRSYSCOM		8b. OFFICE SYMBOL (If Applicable) PMA-240R	9. PROCUREMENT INSTRUMENT IDENTIFICATION NUMBER	
8c. ADDRESS (city, state, and ZIP code) NAVAIRSYSCOM, WASHINGTON, D.C. 20361-1240			10. SOURCE OF FUNDING NUMBERS	
			PROGRAM ELEMENT NO.	PROJECT NO.
			TASK NO.	WORK UNIT ACCESSION NO.
11. TITLE (Include Security Classification) Training Methodologies for Dependent Speech Recognition (SR) Systems (Unclassified)				
12. PERSONAL AUTHOR(S) CDR Richard L. Miller				
13a. TYPE OF REPORT Master's Thesis		13b. TIME COVERED FROM OCT 89 TO MAR 91	14. DATE OF REPORT (year, month, day) March 1991	15. PAGE COUNT 32
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	GROUP	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 a dependent SR system would perform with different accuracies given different ways in which it was trained. The experiment used a SR system (Voice Navigator) which is based on Dragon Systems, Inc. (proprietary) technology. Fifteen subjects trained three different voice patterns each and conducted four tests to compile statistics about the recognition accuracy for each pattern. The experiment was successful and demonstrated that the training methodology used can have significant impact on the performance of a dependent SR system. 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 Richard L. Miller			22b. TELEPHONE (Include Area Code) 408-646-2174	22c. OFFICE SYMBOL Code 37

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.

Training Methodologies for Dependent
Speech Recognition (SR) Systems
by

Richard L. Miller
Commander, United States Navy
B.S., United States Naval Academy, 1974

Submitted in partial fulfillment of the requirements
for the degree of

MASTER OF SCIENCE IN INFORMATION SYSTEMS

from the

NAVAL POSTGRADUATE SCHOOL
March 1991

Author:

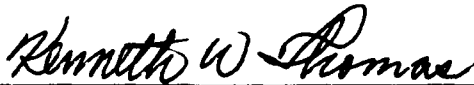


Richard Lee Miller

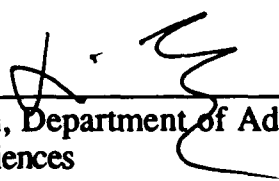
Approved by:



Gary K. Poock, Thesis Advisor



Kenneth W. Thomas, Second Reader



David R. Whipple, Chairman, Department of Administrative
Sciences

ABSTRACT

A research experiment was conducted to determine whether a dependent SR system would perform with different accuracies given different ways in which it was trained. The experiment used a SR system (Voice Navigator) which is based on Dragon Systems, Inc. (proprietary) technology. Fifteen subjects trained three different voice patterns each and conducted four tests to compile statistics about the recognition accuracy for each pattern.

The experiment was successful and demonstrated that the training methodology used can have significant impact on the performance of a dependent SR system. 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 GRAB	☒
DTIC TAB	☐
Unannounced	☐
Justification	
By	
Distribution/	
Availability Codes	
Avail and/or	
Dist	Special
A-1	

TABLE OF CONTENTS

I. INTRODUCTION.....	1
A. BACKGROUND.....	1
B. PROBLEM.....	2
C. SCOPE OF THE THESIS.....	3
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	9
a. 'Subject' Variable	9
b. 'Trial' Variable	10
c. 'Pattern' Variable	15
B. DISCUSSION.....	17
IV. CONCLUSIONS.....	21
REFERENCES.....	23
APPENDIX A	24
INITIAL DISTRIBUTION LIST.....	25

LIST OF TABLES

TABLE I	ANALYSIS OF VARIANCE SUMMARY TABLE.....	10
TABLE II	DUNCAN RANGE TEST RESULTS FOR TRIALS	12
TABLE III	DUNCAN RANGE TEST RESULTS FOR PATTERN	16

LIST OF FIGURES

Figure 1	Average Effect of Trials on Performance.....	11
Figure 2	Individual Effect of Trials on Performance.....	13
Figure 3	Percent Accuracy vs. Trials vs. Pattern	14
Figure 4	Effect of Pattern on Performance	18
Figure 5	Percent Accuracy vs. Pattern.....	19

I. INTRODUCTION

A research experiment was conducted to determine whether a dependent SR system would perform with different accuracies given different ways in which it was trained. The experiment used a SR system based on Dragon Systems, Inc. (proprietary) technology. Fifteen subjects trained three different voice patterns each and conducted four separate trials to test the SR's voice recognition accuracy. Statistics were compiled on each pattern's performance. 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

At present there are many successfully implemented SR systems in the world of business, medicine, assistance for people with disabilities, etc. Most of these systems are of the 'dependent' type, meaning they rely on a speaker to train the SR system to his/her individual voice, i.e. the speaker trains the system by giving the system samples of the user's voice. The system then performs to a certain level of accuracy based on how well it recognizes the voice patterns it was trained with. A dependent SR system's performance depends on how well it can match speech templates with the actual speech characteristics later spoken for recognition. How well a SR system accomplishes this matching depends on the type of algorithm used.

Literature abounds with discussions of how to design algorithms (Lea, 1980; Dixon and Martin, 1979; Waibel and Lee, 1990), however once designed there is

little testing done to determine the best way to train the system for optimum results. Very little can be found in the literature (Lea, 1980; Dixon, Martin, 1979; Waibel and Lee, 1990) regarding proper techniques for training a dependent SR system. Even less is written about differing training methodologies that could possibly be used to optimize SR system performance.

Individual SR systems seem to have 'personalities.' Some perform best when words are spoken relatively fast, others when enunciation is crisp, and still others when words are spoken relatively slowly. The key problem with this uncertainty is the end-user not being provided adequate information to effectively train a particular system for optimum performance. Each vendor addresses the training issue in a general manner, with little or no guidance to the user for optimizing the system's performance.

B. PROBLEM

How do you best train a dependent SR system? The best determination from the literature is to train it in as 'natural' a manner as possible (Lea, 1980; Waibel and Lee, 1990). What is 'natural' to one person is not so to another. Each person has distinctive characteristics about their speech, which is why it is relatively easy for humans to recognize a particular person by the sound of their voice. However, it is more difficult to recognize and identify a particular person's voice if heard over an electronic medium such as the telephone or a radio. The potential for misrecognition increases over such mediums. Such is the problem for a dependent SR system.

A dependent SR system is required to do the very thing which humans have more difficulty doing---matching a specific speaker's voice characteristics via electronic means in order to identify the speaker and accurately interpret the

words that are spoken. In the process of training a SR system, the characteristics of a person's voice are transcribed (via an algorithm) electronically to form a voice template. A SR system's voice templates are created with flaws and artificialities inherent in the tradeoffs associated with choices between algorithms. Therefore, a dependent SR system's recognition accuracy is directly related to the type of algorithm employed, and whether the speaker trains (creates) the templates in a way which optimizes the algorithm's capabilities. Given a specific algorithm, how much impact does the training method have on recognition accuracy? This thesis explores that question as it applies to one specific type of 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 three different training methodologies, utilizing a specific, dependent SR system.

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

Fifteen subjects (six female, nine male) were recruited from the Naval Postgraduate School in Monterey, California. They were all military personnel from the navy and the army. Their ages ranged from 28 to 38. Some subjects had educational knowledge of SR systems, but no one had actual experience using a SR system before this experiment.

B. SR SYSTEM

The SR system chosen was an off-the-shelf product called 'Voice Navigator' by Articulate Systems, which is based on Dragon Systems, Inc.'s SR technology. The algorithm used in the Dragon speech drivers is proprietary. A Macintosh IICx personal computer was used to conduct the experiment. The SR system allows manipulation of three parameters: rejection threshold, number of training passes, and speech input level. The *rejection threshold* can be set on a scale of 0-100% 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 threshold. The threshold was set at 75%, per vendor recommendation, for this experiment (e.g. if the SR system's algorithm determined there was a 75%, or better, chance of matching an utterance with a word stored on the training template, it would display the word). The number of *training passes* allows the user to select how many times a word will be repeated during the training session. Literature indicates that training a word with three to five repetitions yields best results (Poock, 1990). Over five repetitions does not contribute

significantly to improving the quality of the voice template. Three (3) repetitions were used for this experiment. *Speech input level* on the chosen system allows a wide range of volume levels. If spoken too quietly or too loudly the system will prompt the speaker to speak more loudly/quietly. The test subjects were allowed to speak at whatever volume level desired, allowing the SR system to correct volume errors as needed.

A noise-cancelling, "boom" microphone mounted on a headset was used for voice input to the system. Well suited to environments where there is a lot of background noise, such as noisy offices, the noise-cancelling feature allows you to speak quietly in loud environments while retaining high quality results.

C. EXPERIMENT DESIGN

Each subject was given instructions on how to train the SR system. A dialog window on the computer's monitor displayed the word being trained and which repetition the speaker was on. The same vocabulary list of 90 words (Appendix A) was used for creating each template. Three voice templates were created for each subject: Pattern #1--'natural'; Pattern #2--'artificial inflection'; and Pattern #3--'rapid-speak' (see the Testing section which follows).

Each subject conducted, on four separate occasions, a series of test runs against their templates. One test run against each template was conducted during each trial session (total of three test runs for each trial; 4 trials x 3 templates = 12 test runs for each subject; total of 12*15 subjects = 180 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, speaking in a natural manner. The order of the vocabulary words was changed for each trial to prevent the speaker from falling into a speech pattern 'rut.' The subjects

were not allowed to view the computer monitor during trial runs (viewing SR system's accuracy would possibly have altered the manner in which the subject was pronouncing words), nor were they aware of which voice template they were speaking against.

D. PROCEDURE

1. Training

The term 'training' in the context of dependent SR systems refers to the process of a person speaking the words (or utterances) to the SR system that he or she wants the system to recognize at some later point in time. The SR system's algorithm analyzes the voice characteristics and stores the spoken utterances as digital patterns (voice templates). For this SR system, the training procedure consisted of pronouncing each word three times into the microphone.

The first training templates (Pattern #1 -- natural) consisted of 90 vocabulary words, repeated three times by each subject, in a 'natural' manner ($90 \times 3 \times 15$ subjects = 4050 utterances). Each subject created their own, unique Pattern #1 template. Pattern #2's templates (artificial inflection) were created in the same manner, each subject speaking with exaggerated upward and downward inflections on two of the three repetitions, and monotone on the third. Pattern #3's templates (rapid-speak) were again created in the same manner, each subject speaking the words as rapidly as intelligibly possible for all repetitions.

During training, each time an utterance is spoken it is compared to the average voice pattern of the previous entries for that utterance. If not similar enough to the average, it is rejected and the speaker prompted to repeat the utterance. Once the SR system has accepted three repetitions of the utterance, it saves a voice template for that utterance in its memory. For this experiment,

there is a unique template for each word in patterns one, two and three. The patterns are then used by the SR system during testing to compare the speaker's utterance against the respective template from the appropriate pattern. Ideally, the utterance during testing matches its counterpart template in memory and the result is a correct response. In cases where the SR system cannot make this match, a nonrecognition (or rejection) occurs. Occasionally, however, the SR system 'thinks' it has matched an utterance with one in memory, but the match is incorrect. This constitutes a misrecognition. Thus, two types of errors are possible: nonrecognitions (or rejections) and misrecognitions (misinterpretations) of an utterance.(Pooch, Martin, Roland, 1983, pp 2-6) The training procedure took 45-60 minutes for each subject to train all three voice patterns.

2. Testing

Approximately two weeks after all subjects had completed creating their templates, actual testing began. The two week delay was imposed to help dissipate any 'bad habits' developed during the training sessions and minimize a particular subject's possible tendency to pronounce words in an attempt to match a particular voice template. The 15 subjects conducted four trials each. Each trial consisted of three test runs (one for each template). A test run consisted of the subject reading through the list of vocabulary words and pronouncing each word one time in a natural, flowing manner. The templates were loaded into the SR system in a random order. The subjects did not know which template was loaded, nor were they allowed to view the monitor during testing. These measures further precluded the possibility that a subject might tailor his or her pronunciation of the vocabulary words in order to increase recognition accuracy

of the SR system (not that any of the subjects had any desire or motivation to do so). These precautions were taken primarily to minimize any subconscious effects on speaking patterns, and to attempt achieving the most consistent *speech* patterns possible 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 and three), *trial* (one through four), and *subjects* (1-15). 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 presents the three-way analysis of variance summary table for recognition accuracy (arc sin transformation of raw data). As evidenced by the F-ratio for each of the variables and combinations thereof, all three variables show a significant effect on results, and there is significant interaction between the variables as well.

2. Impact of Variables

a. *'Subject' Variable*

Some subjects did have an interactive effect with 'pattern' on the SR system's recognition accuracy, meaning some subjects performed better on certain patterns, and other subjects vice versa. As in most experiments, one would expect subjects to differ and this was no exception; however their variance is isolated in this design.

TABLE I
ANALYSIS OF VARIANCE SUMMARY TABLE

Source	df	SS	MS	F-ratio	Prob
Pattern	2	6.16653	3.08327	14.44	<.001
Trial	3	0.317714	0.105905	1.88	0.0312
Subj	14	8.16656	0.583325	17.1	<.001
Pattn,Trial	6	0.425802	0.070967	2.07	0.0648
Pattn,Subj	28	5.97910	0.213539	6.24	<.001
Trial,Subj	42	2.39650	0.057060	1.67	0.0238
Error	84	2.87376	0.034211		
Total	179	26.3260			

b. 'Trial' Variable

The 'trial' variable had individual as well as interactive effects on the results. The individual impact is depicted in Figure 1. On *average*, there is a slightly upward trend in performance as the subjects proceeded from the first to the fourth trial.

To further isolate and analyze the 'trial' variable, Duncan's Multiple-Range test was conducted. The purpose of a multiple-range test involves "...a stairstep approach to the making of multiple comparisons. Instead of making all comparisons in relation to a single critical difference (as in the *t-test*), the size of the critical difference is adjusted depending upon whether the two means being compared are adjacent, or whether one or more other means

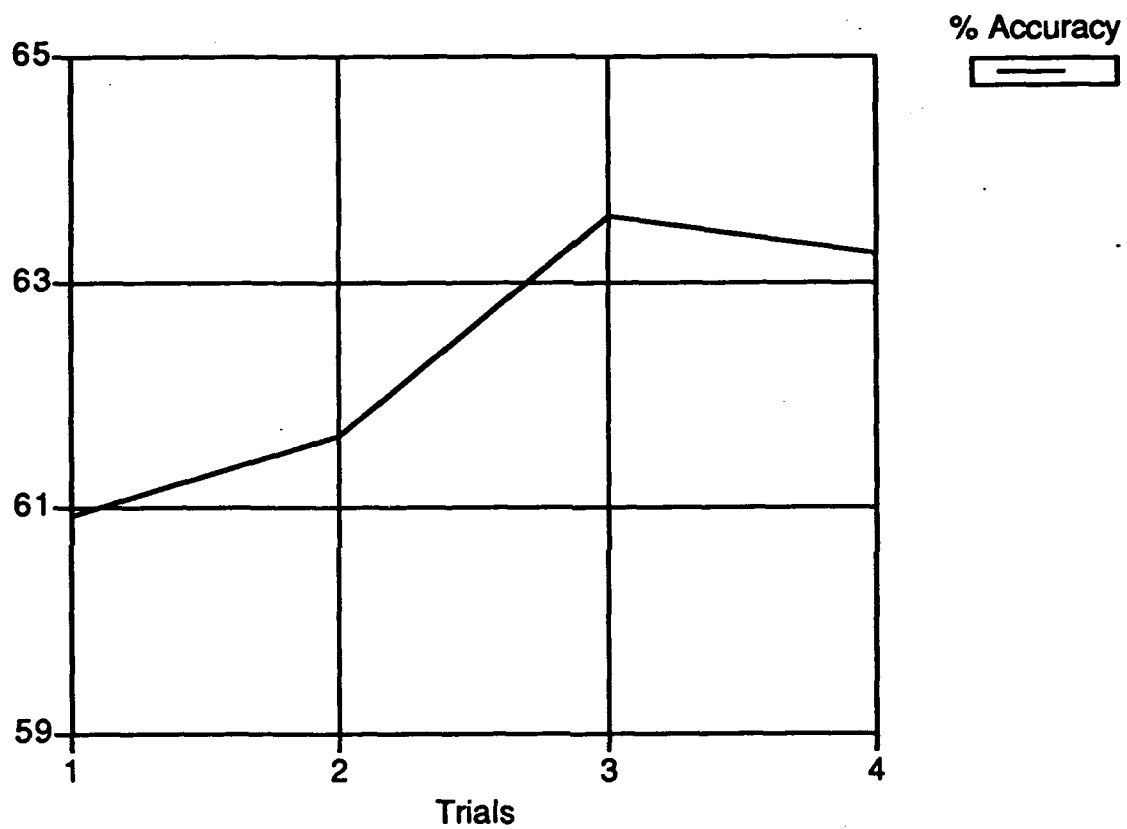


Figure 1
Average Effect of Trials on Performance

fall between those being compared.” (Bruning and Kintz, 1977, p. 116) As seen from the results summarized in TABLE II, performance was significantly affected by the ‘trial’ variable. However, Figure 2 shows this effect is due mainly to the impact pattern three (rapid-speak) trials had on the average.

TABLE II
DUNCAN RANGE TEST RESULTS FOR TRIALS

Rank	Means	r	k	Cdfff Rng	T1 vs.	Effect
T1	2.2771					
T2	2.2918	2	2.77	0.0235	0.0148	Nonsignif.
T4	2.3461	3	2.92	0.0248	0.0691	Significant
T3	2.3817	4	3.02	0.0256	0.1047	Significant
					T2 vs.	
T3		2	2.77	0.0235	0.0899	Significant
T4		3	2.92	0.0248	0.0543	Significant
					T4 vs.	
T3		2	2.77	0.0235	0.0356	Significant

Figure 3 depicts some interesting results regarding the interactive effects between ‘pattern’ and ‘trials’. The performance accuracy for pattern one and two templates is reasonably consistent over all trials. The pattern three

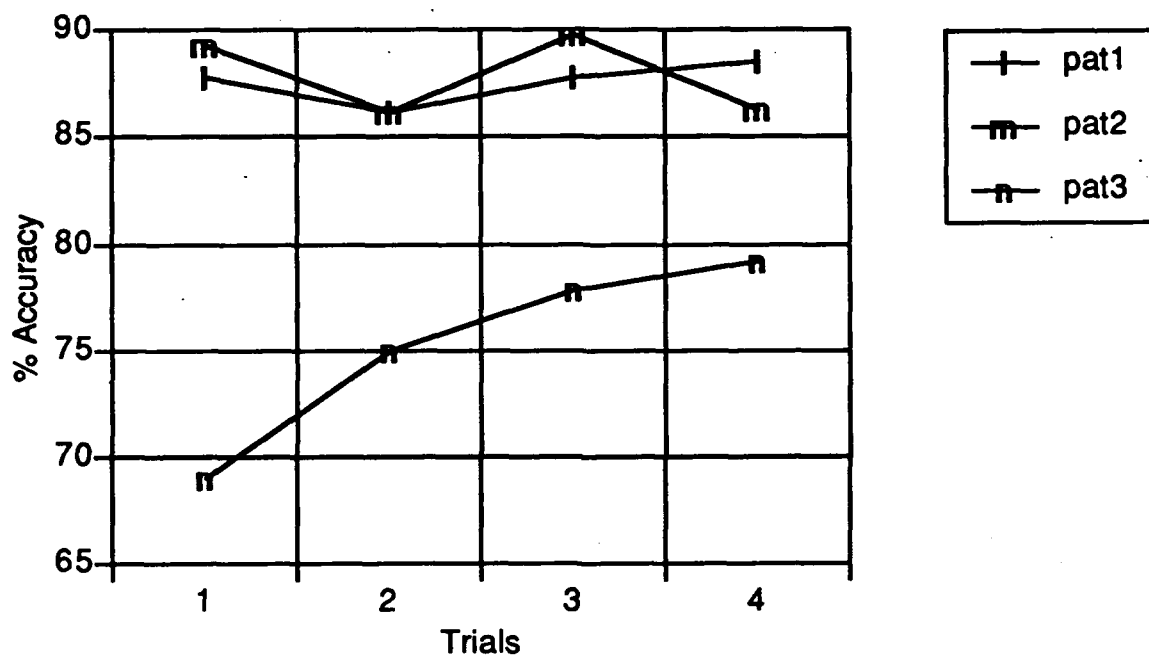


Figure 2
Individual Effect of Trials on Performance

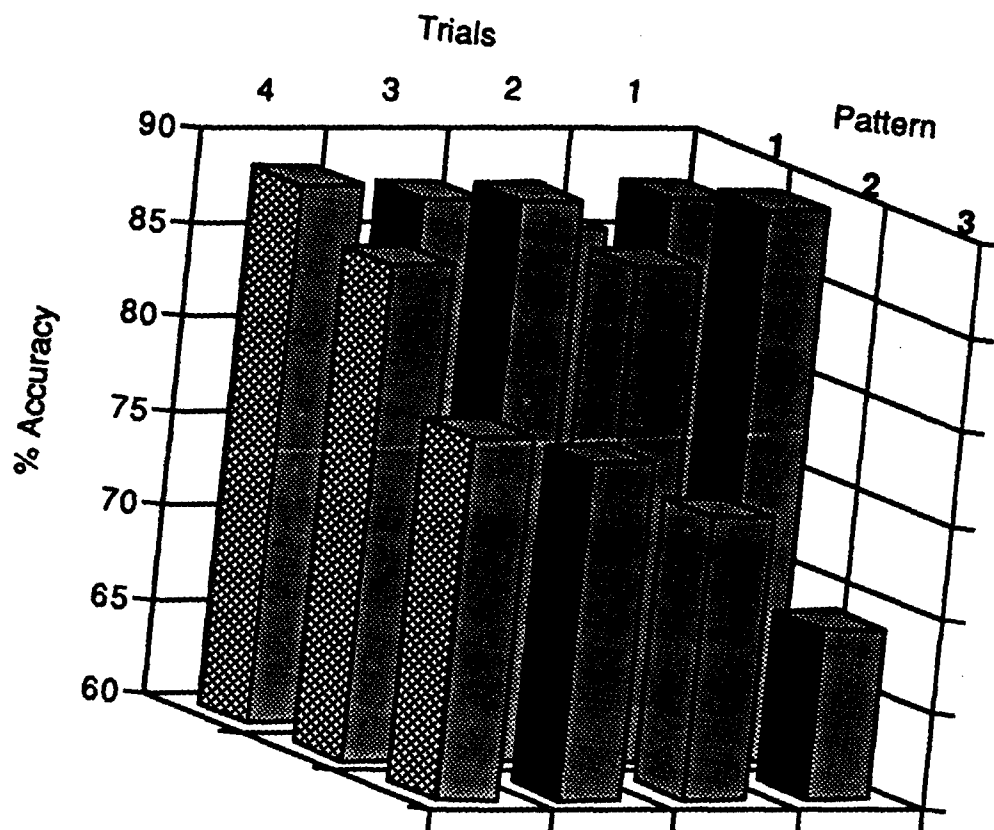


Figure 3
Percent Accuracy vs. Trials vs. Pattern

templates appear to yield much poorer accuracy overall, however the individual effect of the 'trial' variable significantly improves pattern three's accuracy from the first to the fourth trial. A possible explanation for this improved performance over repeated trials would be that speakers become more comfortable 'talking to a machine' (speaking into a microphone and pronouncing words in a more natural manner). Although the 'trial' variable has significant effect on the aggregated performance, in reality it only affects pattern three in a significant manner. This indicates that the methodologies used to train patterns one and two yield consistent performance, independent of a 'learning curve'. From the limited number of trials in this experiment it cannot be determined where the 'flat of the curve' is for pattern three, however it appears to be flattening out between trials three and four, and would probably remain approximately 8--10 percentage points below the performance level of the other two patterns.

c. 'Pattern' Variable

The 'pattern' variable has a significant effect on performance, as depicted in Figures 3 and 4. Figure 4 shows an obvious drop in performance for pattern three on all four trials. To further isolate and analyze the 'pattern' variable, Duncan's Multiple-Range test was conducted. (Bruning and Kintz, 1977, p. 116) The results of the test are summarized in TABLE III. The actual difference of pattern three's results is outside the acceptable range, further supporting the conclusion that the 'pattern' variable has a statistically significant impact on performance results. Of note, the difference between

TABLE III
DUNCAN RANGE TEST RESULTS FOR PATTERN

Rank	Means	r	k	Cdiff Range	P3 vs.	Effect
P3	2.063					
P1	2.436	2	2.77	0.1778	0.373	Significant
P2	2.474	3	2.92	0.1874	0.411	Significant

patterns one and two was .038, less than the acceptable range of .1874, indicating that patterns one and two did not differ significantly in their impact on system performance.

B. DISCUSSION

This experiment did not evaluate whether the overall SR accuracy achieved in the best two examples (patterns one and two) could be improved upon. The recommendation in the SR system's documentation was to train the system in a 'natural' manner, and this was done for one of the training patterns. Pattern two was a variation on the 'natural' theme by attempting to introduce a more dynamic voice pattern with some prosodics, possibly more reflective of the way peoples' voice patterns vary under different circumstances. From the nearly identical results obtained from patterns one and two, it could be asserted that the mean accuracy rates of 87.6 and 87.9 percent, respectively, are as good as this particular SR system might achieve, given the set of vocabulary words chosen for this experiment (Appendix A).

This experiment did demonstrate, in a convincing manner, the downward side of performance using pattern three (rapid-speak). Figures 3 and 4 evidence the poor performance resultant from pattern three. Not only is the performance poor, but the consistency of performance is extremely erratic. The consistency problems resultant from training this SR system in a fast manner are perhaps even more significant than the accuracy issue.

Figure 5 graphically shows the inconsistency of pattern three's performance. Note the consistent performance from patterns one and two (with the exception of a couple of outliers). Additionally, note the performance levels of the four bottom cases from pattern three. These four trials were all

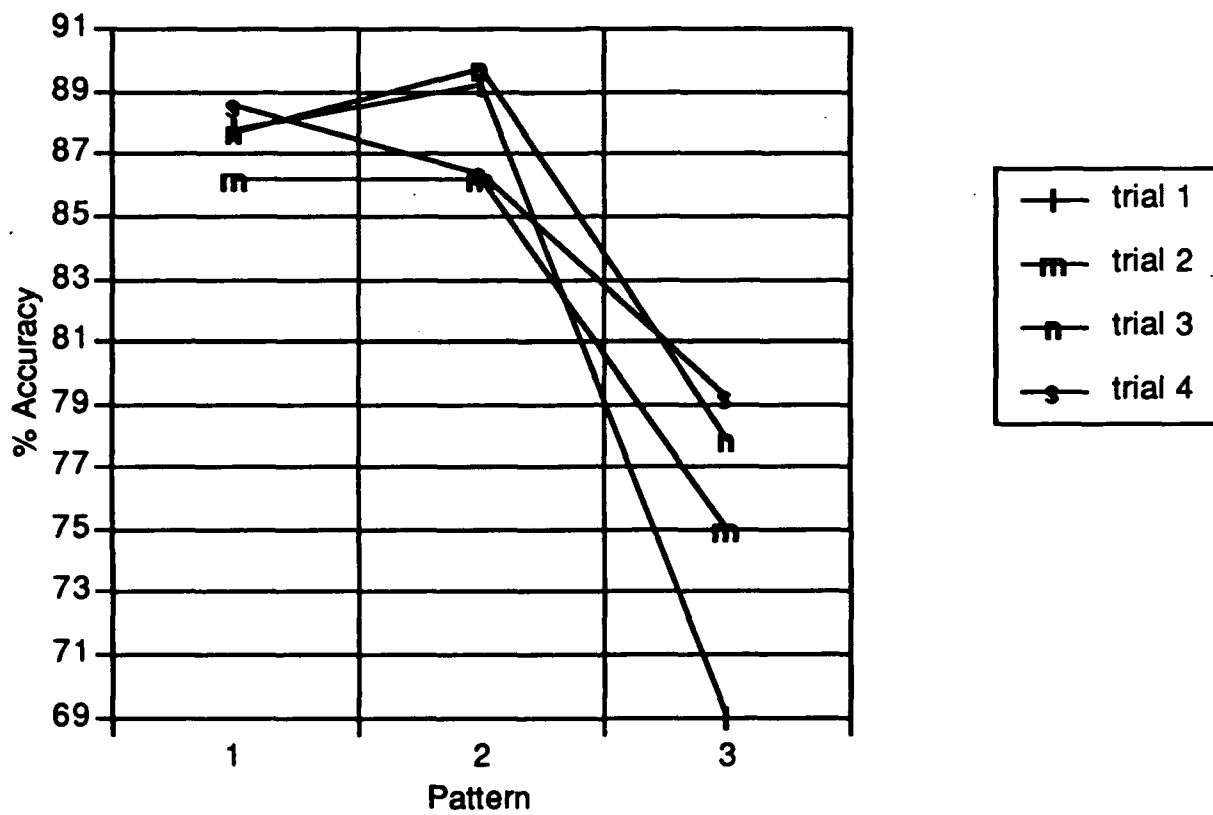


Figure 4
Effect of Pattern on Performance

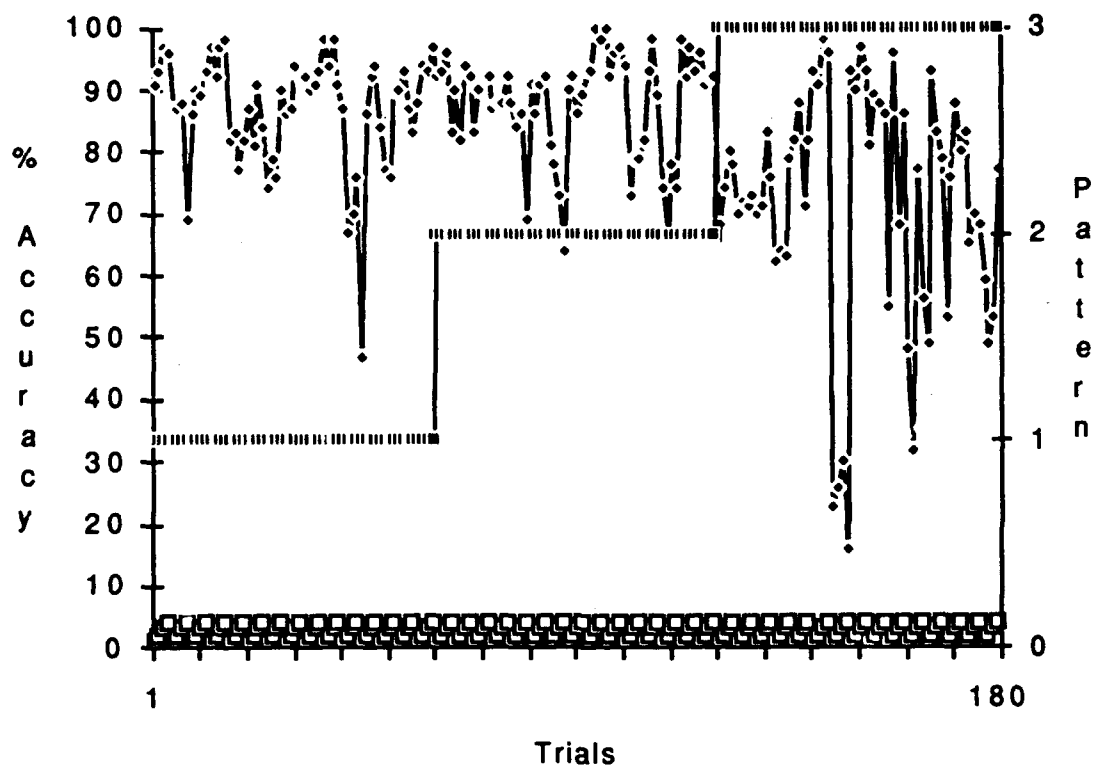


Figure 5
Percent Accuracy vs. Pattern

from the same individual, showing what can happen in the extreme when an individual 'mistrains' the SR system, or for some reason the system performs poorly. From the end-user's perspective, consistency is every bit as important as accuracy, if not more so on many jobs.

IV. CONCLUSIONS

To summarize, the number of trials appears to have an effect only when the voice template was formed under the pattern three methodology. Subjects, as mentioned before, were expected to impact performance, but their variance was isolated for this experiment's design. The effect of pattern, or how the dependent SR system is trained, significantly impacted performance of the system.

In this experiment, patterns one and two did not result in statistically significant performance differences, even though the training methodologies were very different. A conclusion could be drawn that the algorithm employed by this particular SR system was 'tolerant' to pattern one and two training methodologies, however pattern three's methodology (rapid speech) is apparently outside the algorithm's parameters. To support this conclusion, however, a like experiment could be conducted on a different SR system which also employs Dragon Systems, Inc.'s algorithmic approach.

A more general conclusion can be drawn with confidence: the method used to train the chosen dependent SR system *does* affect the recognition accuracy of the system. Patterns one and two resulted in the SR system achieving significantly better, more consistent recognition accuracy than did pattern three. The statistical analysis demonstrates with a high degree of certainty that you can, by accident or by design, train a dependent SR system in an *incorrect* manner, resulting in suboptimal performance. If a person is not given any instructions on how to train a dependent SR system, that person might create voice templates

in a manner which results in extremely poor recognition performance. The user would lose confidence in the SR system's capabilities and most likely avoid using it (particularly if the system is used for a critical requirement).

Manufacturers give little mention of how to train their particular SR systems for optimal results, nor do they suggest alternate methods of training to accomplish that end. A simple statement in the system's documentation such as "...speak naturally...." (which was the case for the system documentation in this experiment) is a catch-all phrase which indicates that the manufacturer may or may not have done any testing to determine the best training methodology to achieve optimal performance.

Even before addressing the issue of how to train a given dependent SR system, a critical question to be answered is what type of algorithm should be designed for the system? This depends on which environment the SR system will be used in (e.g. high stress situations where people's voice patterns vary to extremes, versus the use of voice to augment word processing functions). A dependent SR system can, and should be designed with its users in mind, and the methodologies for training different systems should probably be different in order to achieve optimal performance on each of them. This experiment highlights the need for more research and experimentation to be done in the area of training methodologies for dependent SR 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

- (Bruning, Kintz 77) Bruning, James L. and Kintz, B. L., *Computational Handbook of Statistics*, 2nd ed., Scott, Foresman and Company, 1977.
- (Dixon, Martin 79) Dixon, N. Rex and Martin, Thomas B., *Automatic Speech and Speaker Recognition*, IEEE Press, 1979.
- (Lea 80) Lea, Wayne A., *Trends in Speech Recognition*, Prentice-Hall, Inc., 1980.
- (Neter, Wasserman 74) Neter, J. and Wasserman, W., *Applied Linear Statistical Models*, Richard D. Irwin, Inc., 1974.
- (Pooch 90) Pooch, Gary K., Class Notes from 'Man/Machine Interface' Class, Naval Postgraduate School, Winter Quarter, 1990.
- (Pooch, Martin and Roland 83) Naval Postgraduate School Report NPS55-83-003, *The Effect of Feedback to Users of Voice Recognition Equipment*, by Pooch, Gary K., Martin, B. Jay, and Roland, E. F., pp 2-6, February 1983.
- (Waibel, Lee 90) Waibel, Alexander and Lee, Kai-Fu, *Readings in Speech Recognition*, Morgan-Kaufmann Publishers, 1990.

APPENDIX A

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	