

FREQUENCY EFFECTS IN THE PRODUCTION AND PERCEPTION OF LONG
VOWELS IN TURKISH

ÖZGE SARIGÜL

BOĞAZİÇİ UNIVERSITY

2011

FREQUENCY EFFECTS IN THE PRODUCTION AND PERCEPTION OF LONG

VOWELS IN TURKISH

Thesis submitted to the

Institute for Graduate Studies in the Social Sciences

in partial fulfillment of the requirements for the degree of

Master of Arts

in

Linguistics

by

Özge Sarıgül

Boğaziçi University

2011

Thesis Abstract

Özge Sarıgül “Frequency Effects in the Production and Perception of Long Vowels in Turkish”

This study aims to understand whether the linguistic experience of Turkish speakers have an effect on their knowledge of the phonology of their language and linguistic processes like production and perception. The representation of lexically specified vowel length in borrowed words is chosen due to its special status in Turkish. This type of length in Turkish is not optional or predictable and gives rise to variation and confusion among speakers.

This study consists of two experiments with nonce words and a pronunciation survey designed to understand two types of frequency effects; i) prototype effect ii) exemplar effect in the processes of production and perception of long vowels in Turkish. In order to uncover the prototypical word with long vowels in Turkish, 1722 words with lexically specified vowel length have been sorted out from the official Turkish Language Dictionary (TDK, 1974) and analyzed in terms of i) the syllable number and structure of words, ii) the vowel of the syllable following the long vowel iii) the consonant preceding or following the long vowel. In order to reveal the exemplar effect, phonological neighborhoodness is used.

Results suggest that there is a correlation between production and perception of long vowels in Turkish and the linguistic experience of the speakers. When both types of frequency effects i.e. frequency of patterns and lexical neighborhoodness are used creating nonce words, the versions with long vowels are favored. When they are used independently lexical neighborhood effect appears to be more powerful than the effect of frequency of patterns.

Tez Özeti

Özge Sarıgül “Türkçe’de Uzun Ünlülerin Üretiminde ve Algılanmasında Sıklık Etkileri”

Bu çalışmanın amacı, Türkçe konuşucularının ses üretimi ve algılaması gibi dil yetilerinde sıklık etkisi olup olmadığını anlamaktır. Bu çalışmada Türkçe’ye yabancı dillerden girmiş sözcüklerdeki uzun ünlüler incelenmiştir. Bu sözcüklerde ünlü uzunluğu tahmin edilebilir veya isteğe bağlı değildir ve Türkçe konuşucuları içinde kullanım farklarına rastlanmaktadır.

Bu çalışma anlamsız kelimelerden oluşan iki deney ve bir söyleniş anketi içermektedir. Uzun ünlülerin üretimi ve algılanmasındaki sıklık etkilerine bakmak üzere iki deney tasarlanmıştır.

Deney sonuçları Türkçe konuşucularının uzun ünlüleri üretiminde ve algılamasında sıklık etkisinin görüldüğünü göstermiştir. Sıklık etkilerinin varolduğu anlamsız sözcüklerde diğer sözcüklere göre daha fazla sayıda uzun ünlü kullanımına rastlanmıştır.

ACKNOWLEDGEMENTS

First of all I want to present my deepest gratitude to my advisors Assoc. Prof. Dr. Mine Nakipoğlu and Prof. Dr. Eser Taylan. I would like to thank Mine Nakipoğlu for introducing me to fascinating subjects and a variety of point of views in Linguistics. She has been the perfect advisor that anyone could ever wish for, being extremely kind, patient and encouraging during my studies. Without her support, high spirits and constant interest, this thesis couldn't be completed. I am also grateful for Prof. Dr. Eser Taylan, who made me curious about phonology and who helped me asking questions. She did not only generously share her data I've used in this thesis but also gave me the insights that cleared my thinking throughout this study. I am very lucky to have them as my role models in all aspects of academia and life.

Assist. Prof. Dr. Markus Pöchtrager, who has been great help with his comments in this study, made me think about very intriguing questions of phonology and linguistics in general. Assoc. Prof. Dr. Aslı Göksel read and corrected different versions of this study and Assoc. Prof. Dr. Cem Bozşahin also helped me with his comments and questions. I have to thank them all for being the committee members in this thesis.

I am very thankful to Assist. Prof. Dr. Kathleen C. Hall who read and commented an earlier version of this study and to Assoc. Prof. Dr. Andrew Wedel who was kind enough to suggest and send me the reading materials I needed to complete this thesis.

I also want to thank Mustafa Kaya who helped me in the statistical calculations of the results of the experiments. I am also grateful to the participants of this study, undergraduate and graduate students in Boğaziçi University and residents of the Etiler rest home.

I also want to thank my friends Ayça, Nesli, Reşat, Ercan, Tamer, Nes, Fatma, Güldane, Elçin; who made me laugh, who listened to me, who talked to me, basically who made me happy during my studies. İlay deserves special thanks for being there for me every time I needed.

Finally I want to express my gratitude my family; my brother being my first friend in my life; my father who let me be who I am and let me speak for myself and my mother who always supported me and made me feel safe and confident. This thesis is dedicated to my family.

CONTENTS

CHAPTER 1: INTRODUCTION.....	1
CHAPTER 2: VOWEL LENGTH IN TURKISH.....	6
2.1. Introduction.....	6
2.2. Accounts of long vowels in Turkish.....	8
2.3. Additional types of long vowels.....	12
2.4. Variation among Turkish speakers.....	14
CHAPTER 3: THEORETICAL FRAMEWORK.....	19
3.1. Introduction.....	19
3.2. Implicit Learning Hypothesis and Speech Production.....	21
3.3. Phonotactic Probabilities and Speech Perception.....	24
3.4. Word Processing and Probabilistic Phonotactics.....	25
3.5. Analogical Models.....	26
3.6. Token Frequency.....	30
3.7. Relation Between Phonological Neighborhood and Frequent Patterns.....	31
3.8. Phonotactic Probabilities and Word Segmentation.....	32
3.9. Distributional Cues and Phonetic Discrimination.....	34
3.10. Distributional Patterns in Production and Perception.....	35
3.11. Conclusion.....	37
CHAPTER 4: FREQUENCY EFFECTS.....	38
4.1. Introduction.....	38
4.2. Prototype Effect.....	39
4.2.1 Statistical Distribution of Long Vowels.....	39
4.3. Exemplar Effect.....	47
CHAPTER 5: TESTING THE PSYCHOLOGICAL REALITY OF FREQUENCY EFFECTS.....	49
5.1. Introduction.....	49
5.2. Test I (Production experiment).....	49
5.2.1. Participants.....	50
5.2.2. Materials.....	50
5.2.3. Procedure.....	53
5.2.4. Predictions.....	55
5.2.5. Results.....	56
5.2.6. Stress-Length Interaction.....	63
5.3. Test II (Perception Experiment).....	64
5.3.1. Participants.....	64

5.3.2. Materials.....	64
5.3.3. Procedure.....	64
5.3.4. Predictions.....	65
5.3.5. Results.....	66
5.4. Variation revisited.....	72
5.4.1. Variation among Speakers.....	72
5.5. Pronunciation Survey.....	75
5.5.1. Participants.....	75
5.5.2. Materials.....	76
5.5.3. Procedure.....	76
5.5.4. Results.....	76
CHAPTER 6: DISCUSSION AND CONCLUSION.....	82
6.1. Introduction.....	82
6.2. Nature of the Implicit Knowledge.....	84
6.3. Gradience.....	87
6.4. Redundancy	89
6.5. Token Frequency.....	90
6.6. The Relationship Between Phonetics and Frequency.....	94
6.7. Conclusion.....	95
APPENDICES.....	98
A. Nonce Experimental Items.....	99
B. The Items in the Pronunciation Survey.....	101
C. Results of the Distributional Analysis.....	102
REFERENCES.....	106

TABLES

1. Variables implemented in Analogical Model of Language Algorithm.....	27
2. Percentages of Selected Linking Morphemes When Varying Bias for [-en-].....	29
3. American English Schwa Deletion: Poststress Vowels Preceding Unstressed Sonorant-Initial Consonants Tend to Delete.....	30
4. Distribution of Vowel Types.....	40
5. Distribution of Syllable Number.....	41
6. Position of the Long Vowel in Bisyllabic Words.....	42
7. Position of Long vowel in trisyllabic words.....	43
8. Distribution of Vowel Sequences.....	44
9. Most Frequent Consonants Following the Long Vowel.....	45
10. Least Frequent Consonants Following the Long Vowel.....	45
11. Most Frequent Consonants Preceding the Long Vowel.....	46
12. Least Frequent Consonant Preceding the Long Vowel.....	46
13. Average Well-formedness Rates in Test 2.....	66
14. Words that are Subject to Unusual Shortening.....	77
15. Words that are Subject to Unusual Lengthening.....	79
16. Predictability of German Consonant Pairs.....	88
17. Token Frequency for Unusual Shortening.....	90
18. Token Frequencies for Unusual Lengthening.....	92

GRAPHS

1. Monomodal vs. Bimodal Distributions.....	34
2. Percentages of the Long Vowels in Test 1.....	57
3. Percentages of Long Vowels with respect to Syllable Number.....	59
4. Percentages of the Long Vowels According to Vowel Types.....	60
5. Average Well-formedness Ratings.....	67
6. Ratings for Long Versions for Bisyllabics and Trisyllabics.....	69
7. Ratings for Short Versions for Bisyllabics and Trisyllabics.....	70
8. Rating for Long Versions for /a/, /i/, /u/.....	71
9. Distribution of Following Vowels for /a/.....	93
10. Distribution of Following Vowels for /i/.....	93
11. Distribution of Following Vowels for /u/.....	94

CHAPTER 1

INTRODUCTION

The aim of this study is twofold:

- i) to investigate the issue of vowel length in Turkish with particular focus on the distributional patterning of long vowels, and
- ii) to discuss what effect the lexical distributional regularities of long vowels and the linguistic experience of Turkish speakers would have on the native speakers' knowledge of the phonology of his/her language in the processes like production and perception of vowel length.

Turkish has both phonetic short and long vowels. The co-existence of these two types of vowels raised the question of whether vowel length is distinctive or predictable in the literature. In Chapter 2 these accounts will be introduced.

1. Turkish has 8 distinctive long vowels, /a:, e:, ɪ:, i:, o:, ö:, u:, ü:/, (Özsoy, 2004)
2. Turkish has 4 distinctive long vowels, /a:, e:, i:, u:/ (Lees, 1961)

As for the third account from another perspective, Nuhbalaoğlu (2010) suggests three types of lexically specified vowel length.

In this thesis we will analyze the set on whose status all three accounts agree, that is the borrowed set of words. This set is analyzed as the lexically specified set by all three accounts. The reason that we have chosen this set is the fact that the speakers of Turkish show variation in this set. The variation brings to mind that the clues that we

are looking for may not be categorical but dynamic and probabilistic. That is one of the reasons why we want to address the issue of variation using the statistical information in the data which is the frequency of the patterns. An analysis with the statistical distribution of the patterns will capture the dynamic nature of the variation. Some examples of the variant forms can be seen in (1) and variation with respect to length will be discussed further in Chapter 2.

(1)	elbise	‘dress’	[elbise]~[elbi:se]
	marul	‘lettuce’	[marul]~[ma:ru]
	alfabe	‘alphabet’	[alfabe]~[alfa:be]
	telif	‘copyright’	[telif]~[te:lif]

We try to understand whether there is a relation of the distributional patterns in this set of words and the linguistic processes of perception and production. This variation among speakers makes the question “How do speakers learn, produce and perceive vowel length in Turkish?” more interesting. This is the main question that we will attempt to find an answer for through this study. Other interesting questions regarding the variation are: Why is there a variation in the pronunciation of some words but a consensus in others? Why do only some novel items create confusion but not all? Do these alternating forms share any characteristics? Along with the issues about the variation in familiar words, what kind of clues do the speakers exploit when they decide on the length of a vowel found in an unfamiliar word is also one of the questions that we want to answer.

Attempts to answer these questions will contribute to the discussion of categories and the nature of the phonological/phonotactic knowledge of speakers. The answer to these questions that we will propose is language use; frequency of patterns and phonological neighborhood.

In Chapter 3 various empirical studies are introduced in the areas of areas such as acquisition, perception, production, processing etc. that show the effect of language use in these aspects of language. These studies foster the idea that linguistic representations and processes are directly influenced by the input, that the speakers and learners of the language are sensitive to the statistical information in the language and finally that the language is not purely categorical. In Chapter 3 an umbrella term “usage-based theories” is used to unify the accounts that take close interest in probabilistic information in the lexicon and the probabilistic behavior of the speakers. These accounts also consider frequency in language use as a prominent factor that affects linguistic processes. Under the light of these theories and studies, we will ask the question whether frequent patterns have an effect on the production, perception and the variation of the vowel length in Turkish.

We summarize the questions that this study attempts to answer as such:

- i. Does language use have an effect on the processes of production and perception of long vowels in Turkish?
- ii. Is there a relation/correlation between the distributional patterns and the production/perception of novel items?
- iii. Does phonological neighborhood have an influence on the representation of the words with long vowels?

- iv. Do regularities in the distribution of vowel length have any influence/significance in the variation observed among speakers?

This study is the first psycholinguistic study about Turkish vowel length and it proposes a unified explanation for the behavior of the speakers in production and perception of vowel length in Turkish based on language use. Production and perception are the most important observable parts of the linguistic processes and psycholinguistic experiments that investigate these areas are indispensable for the discussion of representations in language. Being an attempt of a psycholinguistic study that investigates the frequency effects in production and perception of long vowels in Turkish for the first time, this study will also contribute to this field

Two experiments with novel items and a pronunciation survey to understand the relation between language use and production, perception and variation regarding vowel length will be conducted. In Chapter 4 we will review two types of frequency effect that we will consider as a part of language use: prototype and exemplar effect. In order to find out the frequent patterns in long vowels in Turkish we have a detailed analysis of the lexicon with an emphasis on following and preceding consonants, following vowels and syllable structure of the words with long vowels. This analysis reveals certain properties of a prototypical word in Turkish with a long vowel. Second effect we will investigate is the exemplar effect, which means the effect of individual tokens as a whole instead of the patterns we can derive from a part of the word like the following or preceding consonant, etc. As mentioned earlier the goal of the study is to investigate how language use affects the production, perception and variation processes in Turkish long vowels. If language use has an effect on these processes we would expect that the more frequent patterns in the

lexicon and the phonological similarity of the words to influence the production and perception of novel words in our experiments as well as in the variation. In order to test these effects we have conducted two experiments with nonce items, which are discussed in Chapter 5. In the first experiment we test 48 nonce items for the frequency effects in production. 40 participants (mean age 20.6) are asked to produce nonce items, which are designed using different types of frequency effects. Secondly, we have an experiment where we tested the perception with well-formedness judgments. We have 20 participants with a mean age of 20.2. Additionally, we address the question of variation having a pronunciation survey. In this survey we want to understand the nature of change as well as variation; therefore, we have two age groups with 20 participants each, with a mean age of 78 and 20. The results of the experiments and pronunciation survey reveal a significant effect of language use. Finally in Chapter 6 we discuss the implications of the results and the limitations of the study. The relationship of the experiments and the results with respect to concepts like implicit knowledge, gradience and redundancy are explored in this chapter. Some drawbacks of the study and ideas regarding a further study in this subject are also introduced in Chapter 6.

CHAPTER 2

VOWEL LENGTH IN TURKISH

2.1 Introduction

Turkish has both long vowels and short vowels. There have been different analyses in the literature regarding the status of vowel length. First two analyses we discuss below ask the question whether vowel length is distinctive in Turkish or not and if yes in which vowels. Analysis of Özsoy (2004) considers vowel length to be distinctive in all 8 vowels in Turkish. Second analysis by Lees (1961) argues for 4 distinctive long vowels instead of 8. Finally Nuhbalaoğlu (2010) proposes three types of lexical length in Turkish in the framework of the Government Phonology.

After we review these accounts we will introduce the variation that we observe among the speakers. This variation leads us to focus on a certain set of vowels; borrowed set, about which, all three accounts reach a consensus saying this set has lexically specified length.

2.2 Accounts of Long Vowels in Turkish

There are two main analyses regarding the distinctiveness of vowel length in Turkish vowel system.

- i) Turkish has 8 distinctive long vowels. (Özsoy, 2004)
- ii) Turkish has 4 distinctive long vowels (Lees, 1961)

i) Özsoy (2004) argues for 8 distinctive long vowels. Turkish has 8 distinctive short vowels /a, e, ɪ, i, o, ö, u, ü/. (Özsoy (2004), Göksel&Kerslake (2005) among others) which contrast with one another as illustrated in the examples in (2).

(2)	kal	[kal]	‘stay’
	kıl	[kıl]	‘hair’
	kol	[kol]	‘arm’
	kul	[kul]	‘slave’
	kel	[kel]	‘bald’
	kil	[kil]	‘clay’
	kül	[kül]	‘ash’
	kör	[kör]	‘blind’

According to the first analysis, the 8 short vowels each have a long phonemic/distinctive counterpart, namely /a:, e:, ɪ:, i:, o:, ö:, u:, ü:/ in Turkish.

(Özsoy, 2004) to illustrate this issue let us look at the following contrastive pairs in (3).

(3)	damat [da:mat] ‘groom’	damak [damak] ‘palate’
	temin [te:min] ‘provide’	temiz [temiz] ‘clean’
	sine [si:ne] ‘bosom’	sinek [sinek] ‘fly’
	sığlık [sɪ:lɪk] ‘shallowness’	sıla [sıla] ‘renuion’
	doğru [do:ru] ‘correct’	doruk [doruk] ‘peak’
	öğren [ö:ren] ‘learn’	ören [ören] ‘ruins of a building’

düğme [dü:me] ‘button’

dümen [dümen] ‘rudder’

tufan [tu:fan] ‘flood’

tufa [tufa] ‘trick (informal)’

Existence of these pairs suggests that Turkish has 16 vowels, 8 being short 8 being long according to Özsoy (2004).

ii) Lees (1961) observed that all long vowels did not behave uniformly in certain morphologically complex environments. One such context for nouns is their inflected form in the third person possessive. In Turkish the third person possessive suffix has two phonological realizations, [-sı] and [-ı]. As seen in the examples (4) when the stem that the possessive suffix attaches ends with a vowel the suffix has the form [-sı] and when the word ends in a consonant the suffix realizes as [-ı].

			Nominative	Possessive
(4)	(a)	kapı ‘door’	[kapı]	[kapı-sı]
	(b)	duvar ‘wall’	[duvar]	[duvar-ı]

As seen in (5a-b) the words that are very similar in the bare form behave differently when the possessive suffix is attached. This difference in behavior leads Lees (1961) that there are two types of vowel length in Turkish.

			Nominative	Possessive
(5)	(a)	dağ ‘mountain’	[da:]	[daı] *[da:sı]
	(b)	eda ‘mien’	[eda:]	[eda:sı] *[edaı]

In the examples (5) we can see that *dağ* [da:] gets the suffix [-ɪ] as the words that end with a consonant. Lees (1961) argues that these words have an underlying consonant at the end, which determines the choice of form of the possessive suffix. He suggests this underlying consonant is [ɣ], evidence for this consonant comes from different dialects of Turkish (rather than the standard Turkish) where [ɣ] is pronounced. This duality in behavior of some long vowels is not only observable in nouns but also verbs. In Turkish while the verbs that end in a consonant gets the passive suffix [-ɪl], the words that are vowel-ending get the passive suffix [-n] as seen in the examples (6)

- | | | | | | |
|-----|------|-------|---------|--------|---------|
| (6) | yaz- | [yaz] | ‘write’ | yaz-ɪl | [yazɪl] |
| | ye- | [ye] | ‘eat’ | ye-n | [yen] |

However the verbs that end in a long vowel get the suffix [-ɪl].

- | | | | | | |
|-----|----|------|--------|------|-------|
| (7) | eğ | [e:] | ‘bend’ | eğil | [eil] |
|-----|----|------|--------|------|-------|

This example also shows that the words that have “ğ” at the end behave as if they are consonant-final words. Lees concludes that “ğ” in Turkish orthography is the signal for the underlying consonant [ɣ] and the preceding vowels are lengthened and then this consonant [ɣ] is deleted.

According to this analysis he argues for two types of vowel length one being the predictable length, which is signaled by an underlying consonant [ɣ], and the other is distinctive vowel length where the long vowels behave as true vowels as the

word *eda*. In this analysis there are only four distinctive long vowels [a:, e:, i:, u:] which are found in borrowed words mostly from Arabic as seen in the examples (8)

(8)	damat [da:mat] ‘groom’	damak [damak] ‘palate’
	temin [te:min] ‘provide’	temiz [temiz] ‘clean’
	sine [si:ne] ‘bosom’	sinek [sinek] ‘fly’
	tufan [tu:fan] ‘flood’	tufa [tufa] ‘trick (informal)’

Long vowels that are signaled with “ğ” are analyzed as predictable as represented in (9).

(9)	dağ	[da:]	‘mountain’
	eğlen	[e:len]	‘have fun’
	sığ	[sɪ:]	‘shallow’
	iğne	[i:ne]	‘needle’
	doğru	[do:ru]	‘correct’
	öğlen	[ö:len]	‘noon’
	uğrak	[u:rak]	‘haunt’
	düğme	[dü:me]	‘button’

Neither of these analyses is free of problems. Although the first analysis considers vowel length as “distinctive” in all vowel types, Özsoy (2004) still makes a distinction between the sources of the vowel length and splits the words with long vowels into two sets; Turkic words and borrowed words. It is argued that the long vowels in the Turkic set have derived because of the loss of a consonantal element

and now “ğ” in orthography is the trace of this consonantal element. This split is necessary to explain the consonantal behavior of some long vowels in final position as in (5). However it contradicts with the first claim that Turkish has 8 short and 8 long vowels because this split suggests that long vowels in (10) are not identical, that they have different sources.

(10)	eda	[eda:]	‘mien’
	dağ	[da:]	‘mountain’

Second analysis by Lees (1961) is also problematic because the different behavior of long vowels in (5) is only visible in the final position. Therefore there is no evidence (other than orthography) that the vowels in the words (11) behave differently.

(11)	kağnı	[ka:nı]	‘ox-cart’
	anı	[a:nı]	‘sudden’

iii) Nuhbalaoğlu (2010) makes another analysis regarding vowel length in Turkish. Nuhbalaoğlu shows that “ğ” does not represent the same structure in all words and this suggests three types of lexical vowel length: *dağ*-type, *merak*-type (in which there is alternation, discussed in 2.3) and *bina*-type. The behavior of the *dağ*-type words is attested to morphology, not phonology in this analysis. *Dağ*-type words are argued to have a non-branching nucleus and *bina* and *merak*-type words have branching nuclei. The difference between *merak* and *bina*-type words lies in the type of the onsets that follow the nucleus according to this analysis.

Although these three analyses have differences they have a uniform explanation for the borrowed word set, where there is no signal for length. As seen above these accounts treat the long vowels in this set as the lexically specified-unpredictable long vowels.

After examining two more sources of vowel length in section 2.3, we will lay out the variation observed in the borrowed words in 2.4 and discuss the importance of the variation in terms of this study.

2.3 Additional Types of Long vowels

There are two more phenomena regarding vowel length in Turkish.

i) Compensatory lengthening

ii) Vowel length alternation in *merak*-type words

i) Compensatory lengthening (h-deletion)

(12)	kahve	[kahve] > [ka:ve]	‘coffee’
	mehmet	[mehmet] > [me:met]	‘male name’
	tö Ahmet	[tö Ahmet] > [tö:met]	‘accusation’
	ıhlamur	[ıhlamur] > [ı:lamur]	‘linden tree’

In speech, [h] in intervocalic position and [h] before some consonants such as labials [v, m] can be deleted in Turkish. Loss of the consonant [h] in these contexts leads the preceding vowel to lengthen (Sezer, 1985; Kornfilt, 1985). This derivational process is optional and observed in fast/careless speech.

ii) Vowel Length Alternation

Another process regarding vowel length is the vowel length alternation observed in certain stem when a suffix with an initial vowel is attached to it. This process is only encountered in the vowels /i, u, a/, in borrowed words. These words are argued to have a long vowel originally in Arabic and are shortened in Turkish when the vowel is situated in a closed syllable. (Göksel&Kerslake, 2005; Özsoy 2004)

(13)	Nominative	Accusative	Ablative	
	hayat [hayat]	[haya:tɪ]	[hayattan]	‘life’
	zaman [zaman]	[zama:nɪ]	[zamandan]	‘time’
	tetkik [tetkik]	[tetki:ki]	[tetkikten]	‘examination’
	hukuk [hukuk]	[huku:ku]	[hukuktan]	‘law’
	zemin [zemin]	[zemi:ni]	[zeminden]	‘floor’

As seen above, vowel length in Turkish is an intriguing issue. We have a class of borrowed words that sometimes behave differently than a group of Turkic words in which vowel length is signaled in orthography. So far we have seen that in Turkish a group of words (borrowed words) include long vowels and the length should be learned in these words. Speakers possess an implicit knowledge regarding the length of the vowel in these words. Orthography does not always reflect the structure of vowel in term of length, however, its effect should not be totally disregarded in the case of vowel length, because as we will see, in some words in which the vowel length is not signaled with a symbol, speakers may have different choices of the length. For the literate people “ğ” in a certain context facilitates the choice of the

length of a vowel. In section 2.4 we will look at the vowel length variation and discuss its importance for our study. The fact that variation is not observed in the words with the signal “ğ” supports our preference of limiting this study to the borrowed word set with long vowels, since the behavior of the speakers towards these two sets is not in the same manner.

2.4 Variation among Turkish Speakers

Native speakers of Turkish are sometimes puzzled by the vowel length phenomenon. That is, they pronounce long vowels short and short vowels long and also experience difficulty deciding on the length of the vowel in novel or infrequent items. The behavior that some Turkish speakers exhibit can be summarized as displaying;

- i. Free Variation
- ii. Unusual Lengthening
- iii. Unusual Shortening
- iv. Variation/confusion in rare/novel items

We determined the direction of the change, that is we labeled the variation as lengthening or shortening according to the *TDK dictionary* (1974) and Ergenç’s (1995) *Dictionary of Spoken Language*. We have taken the forms in the dictionaries as the starting point and if the pronunciation of the speakers is different from the one in the dictionary we marked the change as the “unusual” form.

- i. Free Variation

Turkish speakers show variation in the use of long vowels. For example, words in (14) show free-variation.

	A	B
(14)	hayır	ha:yır ‘no’
	yarın	ya:rın ‘tomorrow’

ii. Unusual Lengthening

Another variation observed among speakers can be called unusual lengthening because a short vowel is unexpectedly lengthened, as seen in (15).

	A.		B	
	<i>standard</i>		<i>lengthened</i>	
(15)	marul	[marul]	>	[ma:rul] ‘lettuce’
	nasip	[nasip]	>	[na:sip] ‘portion’
	bayan	[bayan]	>	[ba:yan] ‘mrs.’
	hakem	[hakem]	>	[ha:kem] ‘referee’
	tuvalet	[tuvalet]	>	[tuva:let] ‘toilet’
	akraba	[akraba:]	>	[akra:ba:] ‘relative’
	alfabe	[alfabe]	>	[alfa:be] ‘alphabet’
	demokrasi	[demokrasi]	>	[demokra:si] ‘democracy’

The forms in (15B) are not part of Standard Turkish, they are considered as unnatural/marginal by most of the speakers. The variation is not a dialectic difference; speakers may lengthen all the forms above or they may lengthen only one of the forms.

iii. *Unusual*¹ Shortening/ Weakening of Vowel Length

In (16), there is a tendency to shorten the long vowels as opposed to the unusual lengthening illustrated in (15).

	A		B	
	standard		shortened	
(16)	akide	[aki:de]	>	[akide] ‘a type of candy’
	elbise	[elbi:se]	>	[elbise] ‘dress’
	telif	[te:lif]	>	[telif] ‘copyright’
	defile	[defi:le]	>	[defile] ‘fashion show’
	endişe	[endi:şe]	>	[endişe] ‘worry’
	gariban	[gari:ban]	>	[gariban] ‘poor’
	aşiret	[aşı:ret]	>	[aşiret] ‘tribe’

According to the TDK dictionary (1974) these words have long vowels, though native judgments vary. In the off-line mini-survey that is conducted with 11 informants, (25-year old university graduates) the participants are asked to state the more natural form for the pairs in (16).² They are also asked to rate the acceptability of the counterpart/non-preferred one. There seems to be a consensus on the forms [elbise], [endişe] and [gariban] with short vowels, as the preferred forms, but [elbi:se] with a long vowel is not totally unacceptable either. However, most of my informants found [gari:ban] and [endi:şe] with a long vowel totally unacceptable. For the rest of the items most of the time both forms were found acceptable, but the

¹ They are unusual from a prescriptive point of view.

² The word list is constructed according to my judgments as a native speaker.

preferred ones varied from person to person. For example, five of the informants preferred [aşıret] over [aşı:ret] and the rest preferred [aşı:ret]. Judgments mostly favor the words with short vowels; hence there seems to be evidence that vowel length is disappearing/weakening in some contexts. This brings to mind the issue of change in representations through time, which is a topic we will attempt to address later with the pronunciation survey with two different age groups.

iv. Variation/confusion in rare/novel items

As the existing items display variation in terms of vowel length, it is reasonable to expect that Turkish speakers will experience problems in the pronunciation of some novel items. The best mean to observe this confusion is the pronunciation of proper names. For instance, through observation I can say, the words Nakipoğlu³, Ergani, Daren and Vaniköy⁴ puzzle some people and cause variation in pronunciation. The alternate forms can be seen in (17);

(17)	nakipoğlu	[nakipo:lu]	[na:kipo:lu]
	ergani	[ergani]	[erga:ni]
	daren	[daren]	[da:ren]
	vaniköy	[vaniköy]	[va:niköy]

One important point to highlight about the variation above is that all the words that alternate are from the borrowed word set, i.e., the words that have long vowels where the vowel length is not signaled via a symbol in orthography. The existence of a symbol for lengthening seems to be eliminating the possibility of variation.

³ It is a compound that can be represented as *nakip oğlu*, the word follows the stress pattern of compounds of Turkish.

⁴ It is also a compound which can be separated as *vani köy*.

Looking at the examples of vowel length variation we may be led to think that vowel length in Turkish is random or in free variation. However this would be misleading. There are many examples that would be the opposite, as illustrated in (18).

Words with non-variant short vowels

(18)	araba	[araba]	*[a:raba] *[ara:ba] *[araba:]	‘car’
	ilaç	[ilaç]	*[i:laç]	‘drug’
	uzun	[uzun]	*[u:zun]	‘tall’

Words with non-variant long vowels

(19)	bazen	[ba:zen]	*[bazen]	‘sometimes’
	lazım	[la:zım]	*[lazım]	‘necessary’
	temin	[te:min]	*[temin]	‘provide’
	ilan	[i:lan]	*[ilan]	‘advert’

In the examples above there is no room for vowel length change. The more detailed analysis of these words will be introduced in Chapter 5, in the pronunciation survey. These examples clearly show that vowel length is not a random process. The existence of both variant and non-variant forms makes the issue of vowel length intriguing. Throughout our study we will try to address the following questions: i) How do speakers learn, produce and perceive vowel length? ii) Why is there a variation in the pronunciation of some words but a consensus on others? iii) Why do only some novel items create confusion but not all? iv) Do these alternating forms share any characteristics?

CHAPTER 3

THEORETICAL FRAMEWORK

3.1 Introduction

The present study explores usage effects in production and perception of long vowels in Turkish. We will consider two types of effects: frequent patterns and the effect of phonological neighbors. In the last decade the classical theory is challenged by accounts that consider usage as a main component that influences language based on studies that employ frequency and probabilities in their analyses (Bybee 1985, 2001, 2006; Bod 2003; Pierrehumbert 2003a; Goldrick&Larson 2008; Treiman et al. 2000; Saffran et al. 1996a-b, among others). Usage-based accounts advocate for dynamic, gradient representations as opposed to the generativist models which argue for categorical, discrete representations, which are independent from linguistic experience and frequency in the data. The idea that mental representations of linguistic items are directly affected by linguistic experience, i.e. usage, is the central argument that distinguishes usage-based approaches from classical approaches. Roots of the usage-based models go back to the studies by Rosch (1973, 1978), which introduced the idea of non-discrete categories. The properties of the mental representations of linguistic items are very similar to the properties of non-linguistic items since they are the product of the same cognitive organ, the brain. This is the assumption that usage-based models make.

Furthermore these models argue that the mental representations are not based on abstract/minimal rules but are the categorizations of existing items even if this may suggest redundancy in representation. Absence of economical abstract rules independent of usage does not mean that there is no room for generalizations in usage-based models; in fact there are different levels of abstraction and generalization, however they emerge from the lexicon or corpus i.e., linguistic experience. Another characteristics of this model is reflected by the term dynamic. The representations are influenced by experience, and the experience has a dynamic nature, hence, the representations are subject to change any time if the change in the experience is drastic enough. (Bybee 2001, 2006; Bybee&McClelland 2005; Langacker 1991; Bod 2003; Pierrehumbert 2001, 2003a, among others).

After this brief introduction, some basic concepts like frequency, language use, implicit learning, gradience, probabilities will be discussed in reference to specific empirical studies. The usage-based accounts prioritize the psychological reality of the concepts they argue for, that is why the psycholinguistic studies in acquisition, learning, production and perception are indispensable aspects of the usage based accounts. How implicit learning is affected by frequency in language use is important because phonotactic knowledge that a speaker possesses is a type of implicit knowledge and how it is acquired and what factors may affect this process is an intriguing issue. Gradience is also a prominent concept that usage-based theories and probabilistic approaches to linguistics address. How the gradience in the input can effect the linguistic processes will be discussed further with the studies.

3.2 Implicit Learning Hypothesis and Speech Production

Speakers possess a kind of implicit knowledge about their language. Phonotactics is one type of implicit knowledge every speaker has. Dell et al. (2000) investigated whether implicit learning process is affected by experience in adults. Dell et al. (2000) explored the effect of linguistic experience on language production. Their premise was that the language processing system learns the patterns in the language by experiencing (hearing and producing) the sound sequences as well as storing them in the memory. In this work, Dell and colleagues tested the relation between implicit learning and experience by conducting three experiments. The properties of the learning mechanism that they propose were: i) it is sensitive to recent experience ii) it is implicit, i.e. there is no overt intention in learning and iii) it can make generalizations. In order to investigate the relation between experience and production Dell et al. resorted to the use of errors in speech. They based their study on two properties of speech errors; that the speech errors are more likely when the alternating sounds are in the same syllable (*leading list* instead of *reading list*, sound in onset position, [r], is replaced by an onset [l]), and this likelihood is strengthened by language-specific constraints, for example for English /h/ is always in onset and /ŋ/ is always in coda position, so we can expect “reng king” instead of “red king” but we almost never have an error as “nged king”. With these two properties of speech errors in mind, Dell et al. (2000) attempted to make participants acquire and produce specific nonce speech stream with different phonotactic constraints. Then they counted the speech errors subjects produced and compared them with the constraints to see whether these constraints were learned and were effective in production. They

asked participants to produce certain speech streams like “feng keg hem nes” and 95 other strings with a metronome, they repeated this session for four days and they counted the tongue slips (mispronunciations) by the participants. The idea was that the speakers would be sensitive to the position of the consonants in the speech stream when they mispronounced. For example if a participant says “keg ken heng fes” instead of “meg ken heng fes” this will be a legal error because the onset is replaced with another onset in the string. ($k > m$) However if the participant replaces an onset with a coda in the data this is an illegal error (for instance: *seg* instead of *meg*).

They had two different sets of constraints. In one condition they used /f/ only in onset position and /s/ in only coda position. (*fes* condition), in the second condition they used /s/ in onset and /f/ in coda. (*sef* condition). They used /m/, /n/, /k/, /g/ evenly distributed in each position; /h/ was always in onset and /ŋ/ was always in coda position. The vowel was always /e/.

They counted the legal and illegal errors in the data and compared these results with the experimental conditions (*fes* and *sef*). If /s/ and /f/ were subject to illegal errors as the other consonants, then this would show that experimental conditions are not learned and implemented by the speakers. However the results showed that they were learned. Only in 2.3 % of instances of /f/ and /s/ were misplaced illegally, however /n/, /m/, /k/ and /g/ were misplaced for 31.8 % of the whole errors. That shows that participants were able to learn the conditions of *fes* and *sef*, that is they learned the positions of the consonants in the speech stream.

In the second experiment, they repeated the first experiment with /g/ and /k/, i.e. with *gek* and *keg* conditions. They again found a significant difference between

the illegal errors of other consonants and the items that were controlled (5.3 % vs. 22.5 %).

In the third experiment they introduced a second-order constraint. In the speech stream they changed the position of consonants /f/ and /s/ according to the vowels. Vowel type was the second-order constraint in this setting. In one group (*fas-sif* condition) participants produced sets like “fas, nag, hang, mak” and “nif, sig, kim, hing” while the second groups of participants were exposed to *saf-fis* condition. This means in the first group /f/ was always in the onset if the vowel is /a/ and /f/ was always in the coda when the vowel is /i/ and the opposite for the /s/. The results of this experiment demonstrated that participants were able to learn the more complex conditions like (*fas-sif*) for example. Although /f/ occurs both in coda and onset, participants were able to distinguish the vowel type differences in the words. In the *fas-sif* condition they had illegal misplacement for /f/ and /s/ only for 9.7% of the errors, for other consonants they had an error rate of 23.2 % which resembles the results of the examples discussed earlier. Thus this third experiment shows that people do not only learn that /f/ is onset, but they also learn the pattern of vowel-consonant sequence.

As a result of this study it is demonstrated that even the four days of exposition to a linguistic data can affect the implicit knowledge of speakers. This study shows the speakers' ability to learn patterns and implement them in production. However in a language the patterns are not always categorical, in other words the constraints such that /x/ is always in onset and /y/ is always in coda are rare. Goldrick&Larson (2008) using a similar experimental setting tested the effect of phonotactic probabilities on production. They changed the probabilities of the

constraints in the language that the participants were exposed to. They also used /f/ and /s/, however this time they changed the probabilities. In one condition /s/ and /f/ are in the onset position for 100% and 0%, respectively. In the second condition 80%-20%, and others follow as 60%-40%, 50%-50%, 40%-50%, 20%-80% and finally 0%-100%. Their results showed that speakers learn these probabilistic phonotactics in the speech stream and they reflect that knowledge on the speech errors they produce. These studies support the view that linguistic processes are influenced by linguistic experience and frequency information in this experience. In particular, the second study by Goldrick&Larson (2008) shows how the gradience in the input is reflected in language processing.

3.3 Phonotactic Probabilities and Speech Perception

The effect of the frequent patterns in perception is also an interesting research question. Treiman et al. (2000) showed that English speakers are sensitive to probabilistic phonotactic patterns using a well-formedness judgment test. They tested the frequency effect of VC sequences in words in acceptability judgments. They compared the well-formedness judgments of high-frequency VC's and low-frequency VC's, expecting that words with high frequency VC's would be rated better. For example in one set they had VC sequences /up/, /ɜk/, /uk/ and /ɜp/, first two being more frequent than the last two VCs. They constructed words using the same consonants in the beginning such as: /rup/, /nɜk/ and /ruk/, /nɜp/. The participants listened to these words and asked to rate them in a 1-7 scale; 1 meaning that this word does not sound like English, and 7 meaning the word could be a actual word in English. The results supported the idea that speakers' are sensitive to the

frequency of rhymes in the well-formedness judgments. The participants rated words with high frequency rime with better rates than the words with low frequency rhyme. This study shows that frequency of patterns, that is gradience in the input is reflected in the well-formedness judgments of English.

3.4 Word Processing and Probabilistic Phonotactics

Another contribution to the idea that statistical information is used in linguistic tasks comes from research on word processing. Vitevitch et al. (1997) demonstrated that probabilistic phonotactics –i.e., the frequency of the segments in a particular position and the frequency of cooccurrence rates of segments- is represented in the memory and used in language processing. They used CVCCVC type of words that would differ in the lexical frequencies of CVC's. They had four types of words, high-high (/fʌltʃʌn/), high-low (/lʌnðʌz/), low-high (/gʌbsaɪk/) and low-low (/ðʌɪbdʒaɪz/); nonce words in parenthesis are representatives for the each set. High-high words rated significantly better than all words sets, high-low and low-high did not reflect a significant difference in ratings, and finally low-low words were rated with lower rates. They conducted a second experiment to confirm the effect of phonotactic probabilities with the processing time measures. They used the same items in an auditory repetition task. In this test participants first listened to the nonce-item and then were asked to repeat the item. They measured the accuracy as well as reaction times. The results were consistent with the first experiment, that is, high-high words had the lowest reaction time results whereas low-low words had the highest reaction time measures. As further evidence high-high words had the highest accuracy rates while low-low words was not repeated successfully. These experiments show that

speakers use the information in their memory regarding phonotactic probabilities in language processing.

3.5 Analogical Models

There are also analogical models in the literature that can account for alternations that are not totally rule-governed. For example Eddington (1996, 2000, 2001) suggested an explanation in the framework of Analogical Modelling of Language for the processes like stress assignment, s-weakening and diphthongization in Spanish. For instance; stress assignment in Spanish is not rule-governed (there is no rule can predict the stress in all Spanish words) and there are three options for stress placement: final, penult and antepenult. Eddington (2000) suggested that the stress pattern of unknown words is predicted through the existing word tokens. When the speaker is trying to understand the place of stress in a word, s/he searches for the similar words in their mental lexicon and applies the stress of the similar words. He had a corpus analysis with 4970 most common words, and he coded the phonetic content and syllable structure of the words using 13 variables. Table 1 shows how the variables are implemented for the words *personal* and *hablaron*.

WORD	STRESS	Variables												
		13	12	11	10	9	8	7	6	5	4	3	2	1
<i>personal</i>	Final	–	–	–	0	p	e	r	s	o	–	n	a	l
<i>hablaron</i>	Penult	6	pt	pt	pt	–	a	–	bl	a	–	r	o	n

Note: 6 indicates third person plural; pt indicates preterit tense. – indicates that a variable does not apply.

Variables:

1. The coda of the word's final syllable, if there is one.
2. The nucleus of the word's final syllable.
3. The onset of the word's final syllable, if there is one.
4. The coda of the penult syllable, if there is one.
5. The nucleus of the word's penult syllable, or 0 if the word is monosyllabic.
6. The onset of the penult syllable, if there is one.
7. The coda of the antepenult syllable, if there is one.
8. The nucleus of the antepenult syllable, or 0 if the word is bisyllabic or monosyllabic.
9. The onset of the antepenult syllable, if there is one.
10. Tense, or 0 if the item is not a verb.
11. Tense, if the item is a verb.
12. Tense, if the item is a verb.
13. The person the verb is conjugated for, if the item is a verb.

Table 1. Variables implemented in Analogical Model of Language Algorithm

(Eddington, 2000; p:99)

With the help of these variables the Analogical Model of Language algorithm predicted the place of the stress with an accuracy of 94%.

Further evidence for analogical models comes from German and Dutch (Krott et al. 2001, Krott et al. 2007). They addressed the problem of the choice of the linking elements between nouns in German and Dutch compounds. For example in Dutch they showed that the choice of linking element, [-s-], [-en-], $\emptyset$ is accounted for with an analogical model rather than a rule-based model. The use of linking elements is illustrated in the examples (20)

- (20) thee+ $\emptyset$ +bus ‘tea box’
- pygmee+en+volk ‘pygmy people’
- tabak+s+rook ‘tobacco smoke’

The morphemes [-s-] and [-en-] are phonologically identical with the plural morphemes in Dutch. However the semantics of the linking elements in the compounds is debatable. They are not always associated with plurality, and plurality is not always conveyed with these linking elements. Krott et al. (2001) argued for tendencies regarding the choice of these linking elements instead of rules. For example it is argued that [-en-] usually comes after the word with a plural interpretation. However (21b) stands as an exception for this rule.

- (21) a) boek+en+kast ‘book case’
 b) boek+Ø+handel ‘book shop’

There are other phonological, morphological and semantic rules that were proposed, however, since there are exceptions for these “rules” Krott et al. (2001) argued that they are tendencies instead of rules. After Krott et al. had shown that the choice of these linking elements is not completely predictable with the rules and they suggested that the choice is based on analogy.

They had three production experiments where they used novel compounds with different sets of right and left constituents of the compound regarding the tendency of the linking element. For example for the linking element [-s-], they had three sets of left constituents (L_1 , L_2 , L_3) where the word in L_2 shows strong bias toward [-s-] as a linking element and L_3 shows strong bias against [-s-] and L_2 is in the middle. The words in these three sets were combined with each other (L_1R_1 , L_1R_2 , L_1R_3 , L_2R_1 , etc). The participants were asked to choose a linking element for these new compounds. The results suggested that mostly the left constituent determines the choice of linking word. The right constituent has shown to have a

minor role. Krott et al. (2001) suggested that this study is an evidence for the effect of existing exemplars since the words themselves determine the choice of linking words. Table 2 depicts the results for first experiment with the linking morpheme [-en-].

Table 2. Percentages of Selected Linking Morphemes When Varying Bias for [-en-]
(Krott et al. 2001, p:59)

Percentages of selected linking morphemes when varying bias for -en-
(Positive, Neutral, and Negative) in the left and right compound position.

Left		Right Position					
		Positive		Neutral		Negative	
Position							
Positive							
	en	94.8	(11.2)	96.4	(6.7)	87.4	(15.3)
	not en	5.2	(11.2)	3.6	(6.7)	12.6	(15.3)
	other	0		0		0	
Neutral							
	en	75.0	(23.7)	81.9	(15.5)	58.3	(26.9)
	not en	25.0	(23.7)	18.1	(15.5)	41.2	(26.9)
	other	0		0		0.5	
Negative							
	en	18.1	(19.1)	18.8	(19.9)	6.0	(7.7)
	not en	81.9	(19.1)	81.2	(19.9)	94.0	(7.7)
	other	0		0		0	

Note. Standard deviations between parentheses.

In the results we can see that left constituent is the prominent factor of choice. When the left element has a positive bias towards the morpheme [-en-] the percentages of selecting [-en-] is higher irrespective of the biases in the right element. The same experiments in German linking words also showed the same effect. (Krott et al., 2007).

3.6 Token Frequency

Another frequency effect we can consider in linguistic processes is the token frequency of the items. Bybee (2003) suggests two contradicting effects of token frequency on phonological and morphological change. First, she argues that frequent words are more sensitive to phonetic change, for example, reduction, i.e. phonetic changes affect the frequent words with a faster rate. Bybee (2003) gives American English schwa deletion before /r/, /l/ and /n/ as an example of this kind of phonetic change. It is suggested in Hooper (1976) and Bybee (2003) that deletion process effects high frequency words faster than the low frequency words.

Table 3. American English Schwa Deletion: Poststress Vowels Preceding Unstressed Sonorant-Initial Consonants Tend to Delete (Hooper 1976)

No Schwa	Syllabic [r]	schwa + [r]
every (492)	memory (91)	mammary (0)
	salary (51)	celery (4)
	summary (21)	summery (0)
evening (149)		evening (0)
(noun)		(verb+ING)

Frequencies per million from Francis and Kucera (1982) are given in parentheses.

Hooper (1976) divides the words in three groups.

- i) every, evening: two-syllable words, nonsyllabic [r]
- ii) memory, salary, summary: words can vary, either two syllables or with a syllabic [r]
- iii) mammary, celery, summery: words with three syllables, no reduction

As shown in the Table 3 with the figures in the parentheses there appears to be a direct correlation between reduction and token frequency. The frequent words are the ones that undergo reduction.

However, as a second effect that works in the opposite direction, she suggests that more frequent items are more resistant to change since they are more strengthened in the memory than the infrequent words. Therefore frequent words actually resist any regularization process in the language. For example Bybee (2003) gives the examples of *weep/wept*, *creep/crept* and *leap/leapt* pairs and states they have a tendency to regularize as *weeped*, *creeped* and *leaped*. However high-frequency words like *keep/kept* and *sleep/slept* are not regularized as *keeped* and *sleeped*.

3.7 Relation Between Phonological Neighborhood and Frequent Patterns

So far we have seen the effects of frequent patterns and individual words in various studies. However these two kinds of information are not totally different patterns found in a word and the word as a whole is directly related, since the patterns are derived from the word itself. Although it is difficult to set apart the effect of these two kinds of information, it is also a very interesting question. Bailey&Hahn (2001) investigated the independent influence of phonotactic probabilities and existing items in the lexicon using the wordlikeness judgments of English speakers in monosyllabic nonce words. They had CVC words in order to test two types of frequency effects. First effect that they considered is the phonotactic probabilities; the conditional probability of CV and VC and CVC was calculated. In contrast to the phonotactic information, the second effect considered was phonological neighborhood effect. In

other words, the nonwords were derived such a way that their phonological closeness to an existing items varied. There were two types of words; the words that differed in one sound from an existing word (near-misses) and the words that differed in two sounds from an existing word (isolates). The aim of the study was to calculate the different influence of these effects in well-formedness test. They used 1-9 scale to rate the well-formedness of the words. The results suggested that these two frequency types affect the well-formedness judgments independently. To put it differently, there are at least two types of effects of the input on a linguistic process like wordlikeness judgments; one effect is derived from phonotactic statistics other is from individual items. For example near-misses were rated with higher rates than the isolates. Additionally the words with higher conditional probabilities for CV, VC, CVC were rated with better rates than the ones with lower probabilities.

3.8 Phonotactic Probabilities and Word Segmentation

Saffran et al. (1996a-b) showed that both 8 month-old infants and adults are sensitive to the distributional cues in continuous speech when they set the word boundaries. They used a speech stream, which consists of nonce words with segments that have different transitional probabilities within words and across words to test whether human beings are capable of extracting and making use of statistical cues in speech to determine the word boundaries. Transitional probability is basically the conditional probability between two phonemes. In the experiment with adults they created an artificial language consisting of four consonants (p, t, b, d) and three vowels (a, i, u). They formed 12 syllables using those phonemes and finally they made up 6 trisyllabic (*babupu*, *bupada*, *dutaba*, *patubi*, *pitabu* and *tutibu*) words.

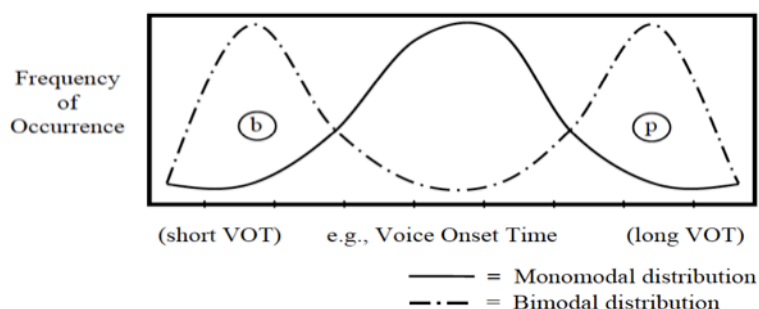
The frequencies of the syllables were not identical to ensure the variability in the data. For example *bu* was encountered for four times in 6 words, while *bi* occurred only once. Because of this variability the transitional probabilities between words and within words were not constant either. Within words, probabilities fluctuated between 0.31 and 1.0 while between words, probabilities did so between 0.2 and 0.1. The subjects were exposed to those 6 words from a speech synthesizer, without any pause or any other prosodic cue such as stress or vowel length. They listened to the speech stream for 21 minutes in a random order in the familiarization phase. In the test phase 6 non-words and 6 part words (eg. *pidata*, *bitaba*) were created. Non-words had zero transitional probability while part words occurred together in the speech stream but their transitional probabilities were lower than the real words. Half of the subjects (n=12) were tested with part words other half tested with non-words. They were introduced with word pairs consisted of one real word and one part word for one group, one non-word one real word for the other and asked to determine which of the words were presented in the familiarization phase. Results showed that subjects successfully segmented words only by depending on the transitional probabilities and they were more successful distinguishing between non-words and real words. They conducted similar experiments with 8-months-old infants using familiarization-preference procedure and showed that infant also detect word boundaries. After infants were exposed to 2-minutes long speech stream they were tested for their listening times. The infants had longer listening times for the non-words and part-words than the real words, meaning they identified real words from the continuous speech stream. Their experiments revealed the fact that both adults and infants are sensitive to statistical information -transitional probabilities- in the

data although at first this kind of information seemed to be too complex to be taken into account.⁵

3.9 Distributional Cues and Phonetic Discrimination

Further evidence for the importance of the distributional cues in acquisition comes from phonetic discrimination. Maye et al. (2002) tested infants (6-8 months old) with continuum of [da]-[ta] sequences to understand how infants differentiate these two phonemes. They used two different stimuli in terms of distribution, one showing a bimodal distribution and the other showing monomodal distribution of the continuum of [da] -[ta] stimuli.

Figure 1: Monomodal vs. Bimodal Distributions



Graph 1. Monomodal vs. Bimodal Distributions

The figure shows the monomodal and the bimodal distribution. If two sounds contrast in a language we should observe a bimodal distribution where we have two peak points in the frequency of occurrence distribution, however if the sounds do not contrast then the distribution should be unimodal, there is only one peak. In the

⁵ This paper is criticized by Boeckx (2010) on the basis of the fact that the artificial language that infants are exposed to is not a realistic model of natural languages. This kind of criticism can be directed to most of the empirical work, however since it is impossible to model the real language acquisition environment in a lab, these kind of artifacts of the designs are present in many of the experiments, not only in linguistics but also in other areas of science.

experiment, the infants of 6 and 8 months were exposed to these two kinds of distribution of [da]-[ta] stimuli and they were found to discriminate the two sounds in the bimodal distribution. The results suggest that 6 and 8 months old infants exploit the statistical cues in the speech when they acquire phonetic categories in their language.

These results are not limited to infants only, Maye (2000) and Maye & Gerken (2000) demonstrate that adults can also learn phonemic categories through distributional information even if there is no minimal pair presented.

3.10 Distributional Patterns in Production and Perception

Zuraw (2000) showed that Tagalog speakers are also sensitive to the distributional patterns in their language. She used an exceptional phenomenon, nasal substitution in Tagalog to illustrate this point. The example of nasal substitution can be seen in (22).

- | | | | | | |
|------|----|----------|------------|----------------|------------------|
| (22) | a) | pighati? | ‘grief’ | pa-mi-mighati? | ‘being in grief’ |
| | b) | po?ok | ‘district’ | pam-po?ok | ‘local’ |

The initial sound /p/ in the first example is replaced with an /m/ when a prefix that ends in a nasal sound is attached to the word. This process is not predictable in Tagalog, for example the same consonant [p] behaves differently in (22 a-b). Although this process seems to be random, Zuraw (2000) showed that these words that alternates share some common patterns and that the speakers are sensitive to these patterns in production and perception. 1736 words that are possible candidates for alternation were analyzed and two main tendencies were found. /p/ and /b/ tend

more to alternate with a nasal than /k/ and /g/; and voiceless obstruents are more likely to be replaced with a nasal. In order to test the psychological reality of these patterns she conducted two experiments with nonce words. In the first experiment she investigated the productivity of nasal substitution and the frequent patterns. Zuraw (2000) tested whether the nonce items that share frequent patterns with the alternating words (subject to nasal substitution) are more likely to undergo nasal substitution. The second experiment made use of the well-formedness judgments to understand the effect of patterns in the perception process of nasal substitution. The participants listened to two versions of the words (substitution/no substitution) and they were asked to rate these words according to their acceptability. The nonce words differed from each other in terms of the degree of the shared patterns with the existing alternating words. In the second experiment Zuraw (2000) observed the effect of lexical patterns. However in the first experiment the subjects used nasal substitution in nonce words with very low rates, which means the frequent lexical patterns in Tagalog are not very productive. This study has shown an exceptional situation in Tagalog, such as nasal substitution applies in a group of words that shares certain patterns, and the speakers are sensitive to these patterns especially in well-formedness test, although the production task has also shown an effect the productivity of this process, the effect of the lexical patterns in production was not very prominent. To summarize this study shows that a process which at first seems like random, may reveal some tendencies if we have a closer look to the data, and the speakers are sensitive to these tendencies.

3.11 Conclusion

Empirical research reviewed above from various areas such as acquisition, perception, production and processing etc. support the idea that language/representations/rules is/are not independent from input, it is not purely categorical and insensitive to statistical information in the input. Under the light of these studies and theories suggested, in this study we will consider frequency effects to explain the production, perception and variation of long vowels in Turkish. Vowel length in Turkish is chosen because the data on variation suggests that speakers do not have categorical judgments (consensus) about the length of the vowels in certain words and this leads to confusion and variation among speakers. Following the usage-based models, we suggest that statistical information in the lexicon i.e. the frequency of patterns/transitional probabilities and the phonological similarity may effect the processes regarding vowel length in Turkish, since they are strengthened in the memory as they are used, and this information is reflected in production and perception. The relation between usage and linguistic processes will be further examined with two experiments and one pronunciation survey, since in a study where the psychological reality of usage is questioned psycholinguistic experiments are indispensable.

CHAPTER 4

FREQUENCY EFFECTS

4.1 Introduction

As stated earlier one of the aims of this study is to understand the role of usage when Turkish speakers and hearers produce and perceive long vowels in Turkish. There are mainly two kinds of frequency effect that have been considered in this study.

- i) prototype effect
- ii) exemplar effect

on the

- i) production
- ii) perception

of long vowels in Turkish.

Different sources of vowel length are discussed in Chapter 2. This study does not look into the compensatory vowel length or alternation of vowel length; it is limited to a borrowed word set where there is no signal for the vowel length. First reason for this choice is the fact that the variation among speakers with respect to vowel length is observed in this set of words. As we have seen in examples (10) through (13) in Chapter 2, variation is observed only in the borrowed word set where there is no orthographic cue about vowel length. Capturing the reasons behind the variation in existing words and also understanding the varying behavior of the

speakers towards the nonce words is the aim of the study, therefore we limited our study to the set where we observe variation. There is also a practical reason regarding the design of the experiments. In the production experiment we had to introduce nonce items to the participants in writing and we could not use a symbol for length; that would of course contradict with the goal of the experiment.⁶ Because of the reasons stated above we limited our study to words with long vowels that are borrowed and that do not include “ğ” as a cue for the vowel length.

4.2 Prototype effect

A vast amount of research suggests distributional properties of lexical items affect linguistic processes. (Bybee 2001, Bod et al. 2003; Dell et al. 2000, Pierrehumbert 2003a, Goldrick&Larson, 2008, Treiman et al. 2000, Saffran et al. 1996a-b, among others)

In order to lay out the distributional properties of words that include long vowels and understand the nature of prototypical words with long vowels in Turkish, we carried out a statistical distributional analysis of words containing long vowels in Turkish. The results of this analysis are later employed in the experiments to create nonce items in order to test the effect of distributional patterns (prototypes) and frequency in production and perception.

4.2.1 Statistical Distribution of Long Vowels

A statistical study on the distribution of words with long vowels was done to find out the most frequent patterns and the prototypical word with long vowels. As mentioned

⁶ Experiments with illiterate people who are not effected by the orthography can be conducted to understand the nature of long vowels that are signaled with “ğ”

earlier data has been limited to borrowed words with lexically specified vowel length, with no orthographic cue for length. The TDK (Türk Dil Kurumu/ Turkish Language Association) dictionary with phonetic transcription (1974) is scanned through to compile the list of lexically specified long vowels.⁷ For the purposes of this study only simplex words have been analyzed. Compounds and morphologically complex words with productive affixes are excluded from the data under study.

The words under investigation are analyzed with respect to three criteria:

- i. Syllable number and syllable structure of words
- ii. The vowel of the syllable following the long vowel
- iii. The consonant preceding or following the long vowel

The data examined consists of 1722 words which are nonnative, borrowed in large proportion from Arabic. These 1722 words contain 1874 long vowels.⁸ These vowels are /a, i, u, e, ü, o/. Table 4 shows the distribution of these long vowels.

Table 4. Distribution of Vowel Types

VOWEL	N	%
a:	<i>1274</i>	68
i:	<i>417</i>	22,2
u:	<i>152</i>	8,2
e:	<i>20</i>	1
o:	<i>5</i>	0,3
ü:	<i>6</i>	0,3
TOTAL	<i>1874</i>	100

⁷ Data on long vowels was sorted out and compiled by Eser Taylan (unpublished manuscript)

⁸ There are words which consist of two long vowels. The discrepancy between the number of vowels and words is due to this fact.

As seen in Table 4, /a:/ is the most frequent long vowel occurring with a rate of 68%. The second mostly encountered long vowel is /i:/ with an occurrence rate of 22%. The third dominant long vowel is /u:/ and it occurs with a rate of 8%. The remaining three vowels /e, o, ü/ have less than 2 % share in the total count.

4.2.1.1 Syllable Number and Structure

In this section distribution of long vowels with respect to syllable structure, in particular the number of syllables is analyzed.

As Table 5 illustrates the majority of the words containing long vowels are trisyllabic (53%), bisyllabic words with a ratio of 32.9% rank second.

Table 5. Distribution of Syllable Number

Syllable number	<i>n</i>	%
1	12	0,7
2	567	32,9
3	915	53,1
4	209	12,1
5	16	0,9
6	3	0,2
TOTAL	1722	100

i) Position of long vowels

Long vowels in Turkish are observed mostly to occur in the penult position of bisyllabic and trisyllabic words. An analysis of words containing long vowels has

revealed that there are only 12 monosyllabic words which contain long vowels in a CV:C template in Turkish.⁹

As for bisyllabic words, which are 567 in number, as Table 6 illustrates, almost 70% have the long vowel in the first syllable, that is, the penult.

Table 6. Position of the Long Vowel in Bisyllabic Words

Position of the long vowel in Bisyllabic words	<i>n</i>	%
PENULT	348	61.4
FINAL	171	30.2
BOTH	48	8.5
TOTAL	567	100

Trisyllabic words display the same tendency in terms of the location of the long vowels. 72% of all trisyllabic words have the long vowel in the penult. And for 67% of all words which are trisyllabic only the penult has the long vowel. This number is significantly high when compared to the other possibilities as displayed in Table 7 below.¹⁰

⁹ The monosyllabic words that have a long vowel to satisfy the minimal word condition are not included in this data, such as, fa: 'a note', do: 'a note' a: 'letter', Inkelas (1995). The monosyllabic words are:

bap	[ba:p]	kar	[ka:r]	yar	[ya:r]
had	[ha:d]	ram	[ra:m]	zat	[za:t]
hal,	[ha:l]	şad	[şa:d]		
kam	[ka:m]	tul	[tu:l]		
ka:p	[ka:p]	yad	[ya:d]		

¹⁰ A detailed analysis of the words with more than three syllables can be found in the Appendix. Since their occurrence rate is low they are not included in the discussion.

Table 7. Position of Long Vowel in Trisyllabic Words

Position of long vowel in trisyllabic words	<i>n</i>	%
PENULT	612	66.9
ANTEPENULT	152	16.6
FINAL	83	9.1
BOTH 1&3	26	2.8
BOTH 1&2	24	2.6
BOTH 2&3	18	2.0
TOTAL	915	100

ii) Distribution of the syllable with a long vowel

Of the 1874 long vowels, only 41 (2%) are found in closed syllables.

iii) Structure of the syllable that follows the long vowel

Another regularity that is observed in the distribution has to do with the structure of the syllable that follows the long vowel. In bisyllabics, the syllable containing the long vowel is followed by a closed syllable (in 77% of the cases); hence a prototypical bisyllabic word looks like;

(23) (C) V: CVC

Now let us look at trisyllabic words. As mentioned earlier, for 67% of the trisyllabics, the long vowel is situated in the penult. In 53% of these words, a closed syllable follows the penult. Hence, the template for a prototypical trisyllabic would be as in (24).

(24) (C) V C (C) V: CVC

4.2.1.2 The analysis of vowel sequences

First, the distribution of the vowels in the syllable that follows the long vowel is analyzed, i.e., the V_1 and V_2 sequences analyzed in V_1CV_2 structures where V_1 is /a:/, /i:/ or /u:/. Results can be found in Table 8.

Table 8. Distribution of Vowel Sequences

A- TOTAL	<i>n</i>	%	İ- TOTAL	<i>n</i>	%	U- TOTAL	<i>n</i>	%
a:Çi	440	44.3	i:Ca	101	48.3	u:Çi	57	51
a:Ce	351	35.3	i:Ce	68	32.5	u:Ca	27	24
a:Ca	139	14.0	i:Çi	38	18.2	u:Ce	26	23
a:Cü	64	6.4	iCü	2	1.0	u:Cü	2	2
TOTAL	994	100	TOTAL	209	100	TOTAL	112	100

The most striking observation one can make about the data is that the most frequent V_1CV_2 sequences are the ones with vowel disharmony. /a:/ (which is a [+back] sound) is followed by [–back] vowels /i, e/ with a rate of 79% which is considerably high. The situation is also similar for /i:/ and /u:/. /i:/ is followed by a /a:/ with a frequency of 42% and /u:/ is followed by a /i:/ with a frequency of 51%. These sequences violate frontness-backness harmony, a property of Turkish phonology. These results not only reveal a distributional regularity about vowel sequences in words with long vowels in Turkish, but also contribute to the peculiar characteristics of long vowels, being situated mostly in disharmonic words.

4.2.1.3 Analysis of the distribution of preceding/following consonants

The frequency of following and preceding consonants of the long vowels were also analyzed.

i) Following consonants

The most frequent consonants following /a:/, /i:/ and /u:/ are given below in Table 9.

Table 9. Most Frequent Consonants Following the Long Vowel

A	<i>n</i>	%	İ	<i>n</i>	%	U	<i>n</i>	%
a:n	137	11.9	i:m	28	12.7	u:r	26	19.1
a:r	133	11.5	i:r	28	12.7	u:n	16	11.8
a:l	110	9.5	i:l	25	11.3	u:d	14	10.3
a:h	94	8.1	i:k	18	8.1	u:l	13	9.6
a:b	80	6.9	i:d	17	7.7	u:b	13	9.6

As seen in the tables above, the sonorants /n,r,l,m/ have a considerably high rate of occurrence after a long vowel in Turkish.

The least frequently occurring consonants are also important since they show the contrast and lead us to generalize a context that long vowels are less likely to be found.

Table 10. Least Frequent Consonants Following the Long Vowel

A	<i>n</i>	%	İ	<i>n</i>	%	U	<i>n</i>	%
a:ş	13	1.1	i:g	2	0.9	u:g	0	0.0
a:g	7	0.6	i:p	2	0.9	u:j	0	0.0
a:p	6	0.5	i:y	2	0.9	u:p	0	0.0
a:ç	4	0.3	i:ç	1	0.5	u:v	0	0.0
a:j	0	0.0	i:j	0	0.0	u:y	0	0.0

ii) Preceding consonants

Also for preceding consonants there is the dominance of the sonorants /m,l,r/. The most frequent ones are;

Table 11. Most Frequent Consonants Preceding the Long Vowel

A	<i>n</i>	%
ma:	115	9.6
la:	99	8.3
ra:	96	8.0
ha:	93	7.8
ka:	91	7.6

U	<i>n</i>	%
ru:	20	13.8
mu:	19	13.1
su:	19	13.1
hu:	14	9.7
tu:	12	8.3

İ	<i>n</i>	%
ri:	49	12.9
ki:	37	9.8
li:	36	9.5
si:	33	8.7
bi:	27	7.1

Least encountered consonants before a long vowel are;

Table 12. Least Frequent Consonant Preceding the Long Vowel

A	<i>n</i>	%
ca:	37	3.1
ga:	23	1.9
pa:	16	1.3
ça:	4	0.3
ja:	1	0.1

U	<i>n</i>	%
fu:	2	1.4
çu:	0	0.0
ju:	0	0.0
pu:	0	0.0
vu:	0	0.0

İ	<i>n</i>	%
yi:	6	1.6
pi:	2	0.5
çi:	0	0.0
gi:	0	0.0
ji:	0	0.0

There is clearly a difference between the most and the least encountered consonants before or after a long vowel.

To summarize the results obtained so far, the prototypical word with a long vowel in Turkish is as following:

- a- is bisyllabic or trisyllabic.
- b- has the long vowel /a/, /u/ or /i/.
- c- has the long vowel in a open syllable.
- d- has the long vowel in the penult.
- e- has the structure CV:CVC, if it is bisyllabic.
- f- has the vowel sequences /a:/-/i/, /i:/-/a/ or /u:/-/i:/.
- g- has the long vowel situated between sonorants.

This distributional analysis of the words with lexically specified long vowels in Turkish reveals the tendencies in the distribution of the long vowels, the surrounding consonants, following vowels and syllable structure of the words. These tendencies have been tested in production and perception experiments to see whether Turkish speakers use/ exploit them productively and whether the hearers use them perceiving the words designed to check their effect, (see Chapter 5 for the experiments).

4.3 Exemplar effect

Another possible usage effect that needs to be tested is the influence of existing exemplars, i.e. specific tokens in the lexicon. Analogical models suggest that a word itself, as a whole, influences the linguistic processes like the choice of stress or the choice of a linking element between nouns in compounds. The rules can not predict the distribution of stress or specific morphemes; however, the analogical properties of words can predict the distribution to some extent in these studies. (Eddington 2001, 2006; Krott et.al., 2001, 2007). It is also shown that phonological neighborhood positively influences well-formedness judgments on nonce words, i.e.

the words that are close phonological neighbors of existing items are rated with better rates in well-formedness judgments. They are treated more word-like than the nonce words that are not phonological neighbors (Bailey & Hahn, 2001). As mentioned in Chapter 2 in Turkish we have an instance of unusual lengthening in the word [hakem] “referee” in the speech of people who pronounce it as [ha:kem]. The existence of the word [ha:kim] “judge” raises the question of whether there is a relation between the representations of these forms, i.e., whether there are neighboring effects. Bybee (1985, 2003) proposes an associative network model for lexical organization and storage. In this model phonetically and semantically similar items are stored together. For example she suggests, the words “send”, “lend”, “trend”, “blend”, “bend” form connections (because they share [ɛnd]) and stored near to each other. This storage is not a list but it is a network, and if one token is activated then others are also activated. In order to see whether this kind of phonological neighborhood influences the production and perception of long vowels in Turkish we will include phonological neighbors in our experiments as well.

In this chapter we reviewed the prototype effect and exemplar effect. Going through the lexicon and some prototypical properties of words with long vowels were abstracted through the help of frequent patterns. In the following chapter we will use these patterns and also phonological neighborhood in two experiments, one on production and the other on perception.

CHAPTER 5

TESTING THE PSYCHOLOGICAL REALITY OF FREQUENCY EFFECTS

5.1 Introduction

In order to discover the relation between usage and production/perception of long vowels in Turkish two experiments with nonce words have been conducted. The aim of the experiments is to see how the distributional patterns of words with lexically specified long vowels in borrowed set and the existing lexical items affect the production and perception of novel items with respect to vowel length.

The first experiment investigates the production process and the second experiment investigates the perception process.

5.2 Test I (Production experiment)

This experiment attempts to address two main questions;

- i) Do speakers of Turkish use the frequent patterns found in lexically specified long vowels when they produce novel words? i.e. Are the speakers more likely to produce long vowels when the nonce words share frequent patterns that are found in the set of words with long vowels in the lexicon?
- ii) Do speakers of Turkish use specific lexical knowledge they possess productively? i.e. Are the speakers of Turkish more likely to produce long

vowels when the nonce words are very similar to the existing items (phonological neighbors)?

5.2.1 Participants

40 undergraduate students at Boğaziçi University have participated in the experiment. (24 females, 16 males) The mean age of the participants is 20.9. The participants report that they are native speakers of Turkish.

5.2.2 Materials

48 nonce words are constructed and implemented in the experiment. The phonotactic properties of Turkish were taken into account in constructing these words. There were 4 groups of words, each set consisting of 12 items (six bisyllabic and six trisyllabic words). The words with each target vowel (/a:, i:, u:/) are equal in number.

A- PRO: Nonce words that include prototypically long vowels (not similar to existing words): There are 12 words in this set. The distributional features of long vowels and words including long vowels in Turkish (4.2.1) are used to construct these words; hence these words will show the prototype effect. Additionally, to make sure that they only show the prototype effect we have avoided using words that are close neighbors of existing items.

i) only bisyllabic (6) and trisyllabic (6) words are used because a word with a long vowel in Turkish is trisyllabic for 53 words out of 100 words and bisyllabic for 33 words out of 100 words. (LANİZ, KİLANI)

ii) The vowels that are predicted to be lengthened were limited to three vowels [a, i, u] because of the high occurrence rates of these vowels as long vowels, 68%, 22 % and 8 %, respectively. (LANİZ, TİLİMA, RUNİF)

iii) Position of the expected/predicted long vowel in the word is penultimate syllable because both in bisyllabic and trisyllabic words it is the prototypical position with a percentage rate of 66 %. (KİMAS)

iv) The bisyllabic words are constructed to fit the skeleton CV:CVC instead of CV:CV, because with the rate of 77% the long vowel in the first syllable of a bisyllabic word in Turkish is followed by a CVC syllable. (LANİZ, KİMAS, RUNİF)

v) For the trisyllabic words we have used the skeleton CVCV:CV since there is not a pattern revealing itself in these words regarding the structure of the last syllable. (KİLANİ, TUMUNİ, TİLİMA)

vi) Vowel sequences are also determined according to the results of the statistical distribution of long vowels. We have used [a:]-[i], [u:]-[i], and [i:]-[a] sequences in the prototypical nonce words.

vii) Finally the following and preceding consonants are determined by the help of the results of the distributional analysis. Sonorants are used as surrounding consonants, except [k] preceding [i:] in KİMAS, since [k] also has a high occurrence rate.

B- EXE: Similar to the existing words with long vowels (not prototypical): There are also 12 bisyllabic and/or trisyllabic words in this set with predicted long [a], [u] and [i] equally distributed in terms of number. In order to have a close neighborhood

effect, only one consonant either in the beginning or in the final position of the word is changed. When the word starts and ends in a vowel, then the first or last consonant is changed. The nature of the change is also controlled. In order to assure similarity between the distorted item and the existing item, only one feature (either voicing, place of articulation or manner of articulation) of the consonant is changed. For instance, the existing item *hata* “mistake” [hata:] is turned into the nonce FATA¹¹. The actual word *akraba* “relative” [akraba:] is converted to AKRADAA as the nonce test item.

In order to avoid the prototype effect, words that do not share the properties of the prototype are chosen. For example words that have long vowels in positions other than penult, or words that do not have [a:]-[i], [u:]-[i], and [i:]-[a] sequences etc. are used. However since most of the existing items display prototype effect, in some words there may be one property that is also used in PRO set. For example, NIBE is derived from [hi:be], in which the vowel is in the penult (a property of the prototypical words with long vowels), or KEBERRU which is derived from [teberru:] with [r] preceding (a property of prototypical words). There are only a few instances that coincide with the prototypical words in EXE set. However in PRO set all the common properties of the prototypical words are used.

C- (NONE): Non-prototypical and non-exemplar: These words are also constructed according to the results of the distributional analysis of long vowels. Least frequent properties have used to create the words in this set.

¹¹ The nonce items that end in [a, e, i, u] may homonyms with the words that are in accusative or dative form. For example *fata* may be interpreted as *fat-DAT*. In order to avoid this confusion we used pictures and told participants the word is the “name” of the picture. Also the nonce items in the sentences are always produced in bare forms.

- i) [a]-[ɪ], [a]-[a], [u]-[u], [u]-[a], [i]-[i] and [i]-[e] sequences are used. Ex:
VAÇAP, ŞIPEZ, ÇUYUK
- ii) Least frequent consonants with long vowels such as [p], [ʃ], [ç] are used as preceding or following consonants.
- iii) The CVCVC and CVCVCV structure is preserved to have consistency between items.

D- (BOTH): Both prototypical and similar to existing words: This set is constructed to see the combined effect of distributional properties and exemplars. For example existing items like [na:rin], [mina:re] are chosen and only one consonant from the end or the beginning is alternated to derive NARÎM and NÎNARE. These words not only show neighborhood effect but also prototype effect since they share many properties with prototypical words with long vowels, such as vowel sequences, preceding and following consonants, position of long vowels.

5.2.3 Procedure

In order to make participants produce nonce items a reading task has been used and to create a natural conversational environment certain meanings are attached to nonce words, for example, a nonce word is said to refer to a special kind of flower, a *new* tool, a color name etc. First a picture with the nonce word is introduced and the participants are expected to learn the meaning of the nonce word. In the second slide, a sentence or a short dialogue is introduced and the participant is asked to read out the entire sentence or the dialogue in which the nonce word would occur. In the sentence the place for the nonce word is left blank and a small icon is used to elicit the nonce word. In this way, we attempt to minimize the effect of ‘listed reading’ as

participants are encouraged to produce the words not by reading but through recall. The sentences are designed in such a way that all nonce items are in the nominative case, which is a null affix in Turkish. In order to avoid any semantic effect the attached meanings (pictures) of the nonce words are controlled. PRO, EXE, NONE and BOTH sets are attached with the same group of pictures such as flowers, spices, colors, tools, birds etc.

To give a sample protocol:

The participants are shown the picture below:



Slide 1. Sample test item

In the next slide the participants are presented with a sentence in which nonce item is represented and the participant reads out the sentence.

BAZILARI KIRMIZI ETİ ____ _BAHARATIYLA
PİŞİRMEYİ SEVER.



(SOME LIKE TO COOK MEAT WITH THE SPICE
_____.)

Slide 2. The reading task

The experimental sessions took place in a soundproof room and the sessions were audio recorded.

In order to familiarize the participants with the task, two examples with the real words were introduced. The procedure of the experiment was presented using these two examples. In the experiment some real words from Turkish are used with the nonce items in order not to lose the attention of the participants.

5.2.4 Predictions

As stated earlier the goal of this experiment is to show whether the speakers of Turkish use general distributional patterns in the lexicon and more specific information presented in individual words regarding vowel length when they produce novel items.

In the present study we have four sets of nonce words in which we implemented different kinds of frequency information. If the speakers are sensitive to these frequency effects and if they use these effects productively then we expect;

- 1- NONE set to be produced with short vowels.
- 2- PRO and EXE sets to be produced with long vowels more than NONE set.
- 3- BOTH set to be produced with long vowels with higher rates than PRO, EXE and NONE sets.
- 4- We do not have a prediction about the relation between PRO and EXE sets; however, if there is a significant difference between the results of the PRO and EXE sets then it will be interesting to see which effect (general patterns or specific existing items) is more dominant in the production of long vowels.

These predictions are summarized below.

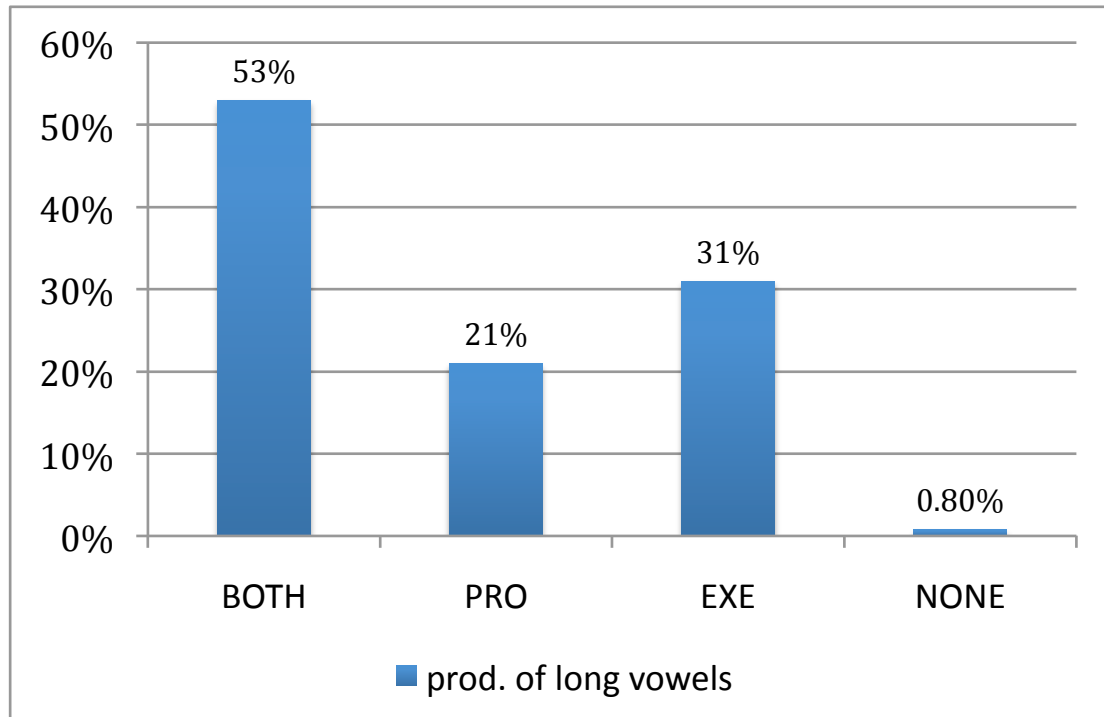
$BOTH_{PRLV} > EXE_{PRLV}$, $PRO_{PRLV} > NONE_{PRLV}$

PRLV: production rate of long vowels

5.2.5 Results

The words that are produced with long vowels were counted in each set. We resorted to the native speaker judgments in order to decide the vowels' length. The percentages were calculated according to the total production rates in each set.

In the set PRO, EXE, BOTH and NONE 21%, 31%, 53% and 0.8% of the nonce words are produced with long vowels, respectively. Graph 2 illustrates the results.



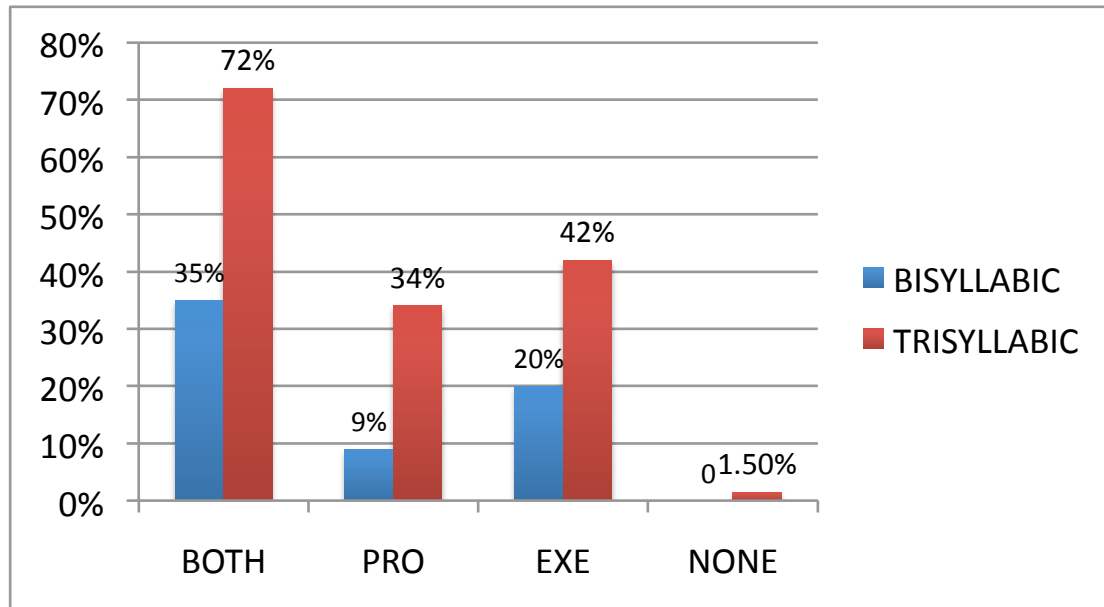
Graph 2. Percentages of the Long Vowels in Test 1

As the results show almost none of the items in NONE set is produced with long vowels. PRO set has a 21% production rate of long vowels in 12 items. This means when frequent patterns found in the words that have lexically specified vowel length is used in nonce words, these patterns increase the tendency of native speakers' production of long vowels as opposed to the infrequent patterns that are used in the NONE set. This result clearly indicates that information about distributional patterns that we have used in PRO words, such as following vowel, following and preceding consonants have an effect on production of long vowels in Turkish. Similarly we have a 31% production rate of long vowels in EXE set. If we compare this to the baseline of 0.8%, we can say that people are more likely to produce long vowels in novel words when the words are similar to the existing items with long vowels in Turkish. These two rates, 21% and 31%, establish the existence of different frequency effects independently, since we have tried to use only one effect in these

two sets. Recall that we have attempted to set these two effects apart as possible and the result obtained actually confirms that we were able to measure two independent effects. When we combine two frequency effects in one set, we obtain a rate of 53% where nonce words are used with long vowels. This is a very neat result conforming that 21% and 31% were really independent and when two types of information (frequent patterns and existing words) are combined, they add to each other and boost the production rate of long vowels. Finally if we compare PRO and EXE set we see that words in EXE set are more likely to be produced with long vowels (31%) than the words in PRO set (21%). This difference is shown to be statistically significant as well. One-way within subjects ANOVA was conducted in order to compare the frequency effects, and all four types of effects in production are found to be significantly different from each other ($F(3,37)=98, p < 0.001$).

Therefore we can conclude that specific words i.e. words as a whole, seem to carry more information about vowel length than frequent patterns. Novel item's similarity to specific existing words increases the possibility of producing a long vowel more than the general frequent patterns that are derived from the lexicon.

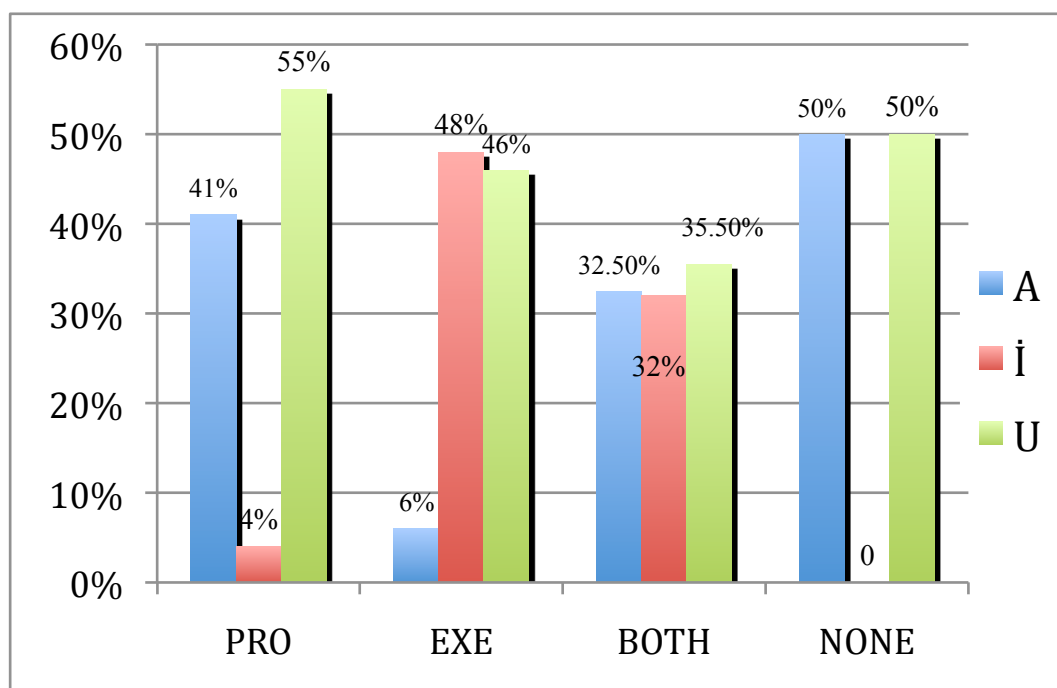
If we look closer into the sets we observe some other interesting factors affecting the production rate of long vowels, like syllable number and vowel type. Having three or two syllables changes the production rate dramatically in all sets. In PRO for example 34 % of the trisyllabic nonce words are produced with a long vowel (penultimate position as intended) as opposed to only 9% of the bisyllabics. Graph 3 shows the percentages for all sets.



Graph 3. Percentages of Long Vowels with respect to Syllable Number

As the numbers suggest in all types of nonce word sets, there is a significant difference between bisyllabic and trisyllabic words. In all sets the production rate of long vowels in trisyllabic words is more than the production rate of long vowels in bisyllabics. These results are consistent with the distributional properties of long vowels in Turkish. Trisyllabics with long vowels outnumber bisyllabics with long vowels in Turkish. Trisyllabic words with long vowels constitute 53% of the all words with long vowels, while bisyllabics constitute only 33% of the total. This discrepancy is reflected in the results of the production test. The differences between syllable numbers are statistically significant as well. ($F(1,39)=102$, $p < 0.001$)

We can also analyze the results according to the vowel types that are subject to lengthening. In Graph 4 the results are summarized with respect to the vowel types.



Graph 4. Percentages of the Long Vowels According to Vowel Types

If we look at the vowel types that are subject to lengthening, there is a discrepancy. /u/'s are lengthened more than /a/'s and /i/'s. In the PRO set there are very few words produced with long /i/'s. For example as we said 29 % of the PRO trisyllabics are lengthened. However, none of these words includes a long /i/. Only /a, u/ are produced as long vowels. This result is not unexpected; although the number of the long /i/ is higher than that of the long /u/'s in Turkish, we observe a tendency towards shortening in long /i/'s. This tendency will be discussed in detailed later in following sections.

This picture changes when we look at EXE set. In this set, bisyllabic words with /a/ are lengthened only once out of 80 utterances/ productions, (i.e., fata:). As for trisyllabic words, the percentage of long /a/ is very low with only 8 instances among 80 utterances. This result is actually unexpected. Long /a/'s constitute 68% of the all the long vowels, so they should be lengthened more than /i/'s or /u/'s if

frequency matters. However, if we look closer to the words FATA and DAMAS which are intended to be the lexical neighbors of ‘hata’ [hata:] and ‘damat’ [da:mat], we see that they are also close lexical neighbors of ‘data, yatak, vatan, yatay’ and ‘dama, damak, damar’ respectively. The words in the second set do not have a long /a/. Existence of these words can actually be the reason for the low rates of vowel lengthening. If we compare these words with SEÇİ, NİBE and DUFAN, which have higher rates of lengthening, we see that these words do not have close lexical neighbors with short vowels. The representation of nonce words may be effected with the coexistence of the words with short vowels and long vowels. This can be an explanation of this situation. The other interesting result is that /u/’s in all sets are more likely to be produced as long compared to /a/’s. This is an unexpected result because of two reasons: First of all, from the frequency perspective, since long /a/’s (68%) exceed the number of long /u/’s (8%), we would expect /a/’s in novel words to be produced as long vowels more than /u/’s. This result is also surprising from a phonetic perspective. It is noted that cross-linguistically long /a/’s are encountered more than long /u/’s. (Lehiste, 1970). This discrepancy can be explained by relative frequency. So far we have been using absolute frequency when referring to frequency of patterns in words with long vowels. In order to talk about relative frequency we need frequency count of the words with short vowels as well. We need to calculate the following ratios in order to have a more accurate comparison;

$$\text{Relative frequency of /a:/ (percentage)} = \frac{\text{Number of /a:/}}{\text{Total \# of a (/a/+/a:/)}} \times 100$$

$$\text{Relative frequency of /u:/ (percentage)} = \frac{\text{Number of /u:/}}{\text{Total \# of u (/u/+u:/)}} \times 100$$

Unfortunately we have not been able to complete the analysis for all the lexical items for Turkish so far. However, we have a partial analysis that shows that if we have relative frequencies of long /u/'s and long /a/'s, the difference will not be as dramatic as 68% and 8%. 10259 bisyllabic words¹² obtained from TDK Manual of Punctuation¹³ are analyzed in terms of their vowels. (Nakipoğlu&Kaya, in preparation) If we look at only bisyllabics we have the following results:

$$\begin{array}{l} \text{Relative frequency of /a:/ (\%)} \\ \text{(bisyllabics)} \end{array} = \frac{785}{6317} \times 100 = 12 \%$$

$$\begin{array}{l} \text{Relative frequency of /u:/ (\%)} \\ \text{(bisyllabics)} \end{array} = \frac{69}{1739} \times 100 = 4 \%$$

Relative frequency of /a:/'s still exceeds /u:/'s according to these results however, the difference between 12 % and 4% is clearly less than the difference between 68% and 8%. This means long /a/'s in Turkish is frequent than long /u/'s but short /a/'s in Turkish are also more frequent than short /u/'s. This result alone still cannot explain the high production rates of long /u/'s. The fact that short /a/'s outnumber short /u/'s in the Turkish lexicon, however, can partly explain why participants insist more on short /a/'s than short /u/'s even when the words possess some prototypical properties of the words with long vowels. Short /a/'s being more prominent than long /a/'s and

¹² These words are not only the simplex words. This group includes derived words as well as the compounds.

¹³ From TDK web site

also short /u/'s may explain this unexpected result. The issue begs for further investigation.

So far we have seen that PRO and EXE sets behave differently in terms of the vowel types. /a/'s tend to be produced with long vowels when they are closer to prototypes than the exemplars and /i/'s do not show any prototype effect. However with EXE words /i/'s tend to be produced long with higher rates compared to /a/'s. In the BOTH set, the numbers are very close to each other. This means all the three vowels are lengthened with the same rate thus we can say the difference in the behavior of the EXE and BOTH sets is neutralized in the BOTH set. Since the BOTH set is designed to have the combined effect of exemplars and prototypes this result is not surprising and also is an evidence for the combined effects.

5.2.6 Stress-length Interaction

Lexical stress the words bear is also taken into consideration in the production experiment since stress and length can be interacting factors. Turkish canonically has stress in word-final position (Sezer 1981, Göksel& Kerslake 2005, among others). Çakır (2000) shows that words that contain a long vowel always have stress in the final position as the words that exhibit a regular stress pattern. In the experiments participants produced the words with final stress whether they used a long vowel or not. This fact rules out the interaction between stress and length in production since all the words tested whether with long or short vowels are produced with a final stress.

5.3 Test II (Perception Test)

The questions for which we seek answers in this experiment are;

- i) Do speakers of Turkish use the information about the frequency of patterns found in lexically specified long vowels when they rate the well-formedness of novel words? i.e., Are novel words that share general distributional information with the words with long vowels rated better on a well-formedness scale when they are produced with long vowels?
- ii) Do speakers of Turkish make use of specific lexical knowledge in the well-formedness judgments about novel items? i.e. Are novel items with long vowels rated with higher rates (more wordlike) when they are phonological neighbors of existing words with long vowels?

5.3.1 Participants

20 undergraduate students with a mean age of 20.2 from Boğaziçi University participated in the experiment. (12 females, 8 males) These participants first attended the production experiment and the second experiment was conducted at least one week after the production experiment.

5.3.2 Materials

Identical nonce items described in the production test were used.

5.3.3 Procedure

A test for well-formedness is used in the second experiment. Participants listened to two versions of the nonce words together, for example [laniz] and [la:niz]. First they are asked to choose the preferable one, and later they are asked to rate the forms from on a well-formedness/ naturalness scale of 1 to 5; 5 being perfectly natural and 1 being not acceptable. In order to make the task less complicated for the participants, they are told to rate the preferable one as 5 and the non-preferable one accordingly, between 1 and 5. This method enables us to capture the probabilistic nature of the speakers' behavior and to make generalizations about tendencies, instead of discrete judgments. Finally the order of the recordings is controlled. 10 participants first listened to the long version and later the short version and remaining 10 listened to the reverse order. The results from the two groups of participants reported together because we could not find an order effect in these two groups. ($F(3,16)=0.7, p>0.5$)

5.3.4 Predictions

If the well-formedness judgments are influenced by frequency effects that are available in the nonce words (prototype and exemplar) we expect BOTH set to be rated with higher rates when they are produced with long vowels and vice versa. For NONE set we expect the versions with short vowels to be rated with higher rates than the versions with long vowels. The main predictions can be summarized as such;

Across Sets

$BOTH_{long} > EXE_{long}, PRO_{long} > NONE_{long}$

$NONE_{short} > EXE_{short}, PRO_{short} > BOTH_{short}$

Within Sets

$BOTH_{long} > BOTH_{short}$

$NONE_{short} > NONE_{long}$

$EXE_{long} > EXE_{short}$

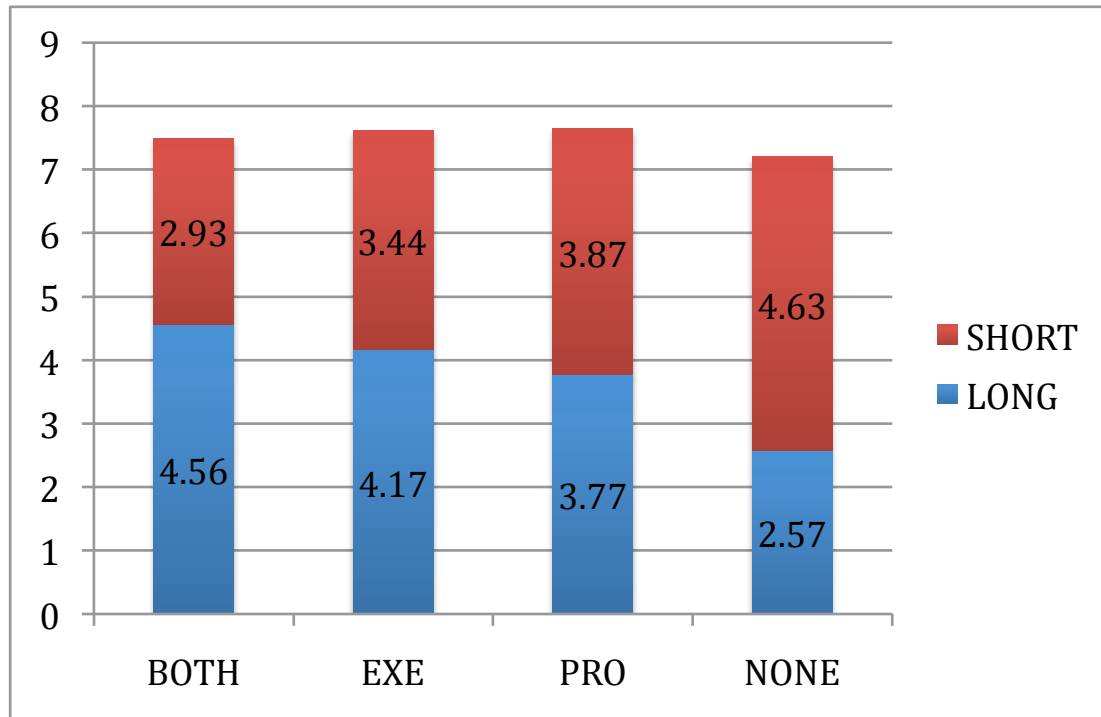
$PRO_{long} > PRO_{short}$

5.3.5 Results

Results are consistent with the predictions. The BOTH set is rated better with long vowels than any other set. And the NONE set is rated better with short vowels than any other set. PRO set confused the participants most; they are observed to be mostly indifferent between two different versions of the words. Well-formedness rates of all versions can be seen in Table 13.

Table 13. Average Well-formedness Rates in Test 2

	BOTH	EXE	PRO	NONE
LONG	4.56	4.17	3.77	2.57
SHORT	2.93	3.44	3.87	4.63



Graph 5: Average Well-formedness Ratings

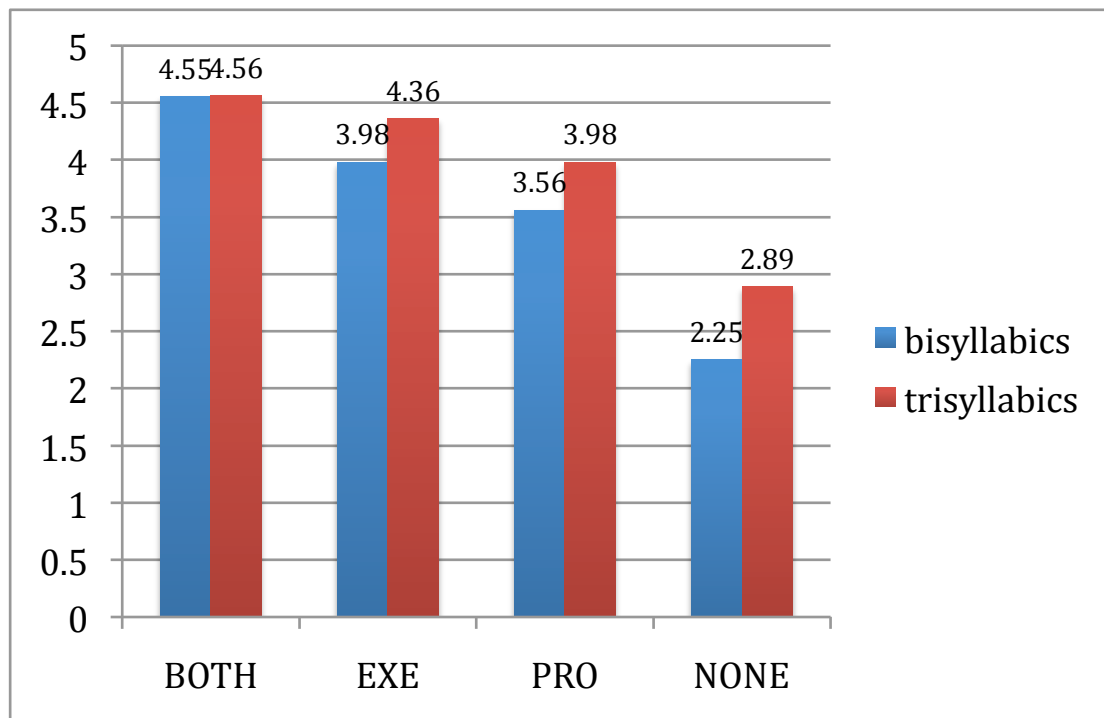
The numbers show the average ratings of the participants for each version of each group. For example, the participants on average rated items in the BOTH set with 4.56/5 when they are produced with long vowels, the short vowel versions of the same set are rated 2.93 on average. As we can see the ratings of the words with long vowels decrease as we move on to BOTH, EXE, PRO and NONE sets. Not surprisingly the ratings for the words with short vowels increase in this order.

The results clearly demonstrate the frequency effects that we are looking for. When a novel word does not share any prototypical property with words with long vowels in Turkish, and when they are not phonological neighbors of these words (NONE) they are rated much better when they are produced with short vowels, 4.63/5 as opposed to when they are produced with long vowels (2.57/5). This difference confirms the fact that novel words are preferable with short vowels when they do not share any general distributional information or specific information with

existing words with long vowels in Turkish. The average rate of 2.57 gives us a baseline, which means that if we add a frequency effect to nonce items we should compare their long vowel rating with 2.57 in order to see the frequency effect. Actually in all sets (BOTH, EXE and PRO) we have higher ratings than this baseline; 4.56, 4.17, 3.77, respectively. This difference in four sets is also found to be statistically significant. One-way within subjects ANOVA was conducted in order to compare the frequency effects, and all four types of effects are found to be significantly different from each other ($F(3,16)=85$, $p < 0.001$).

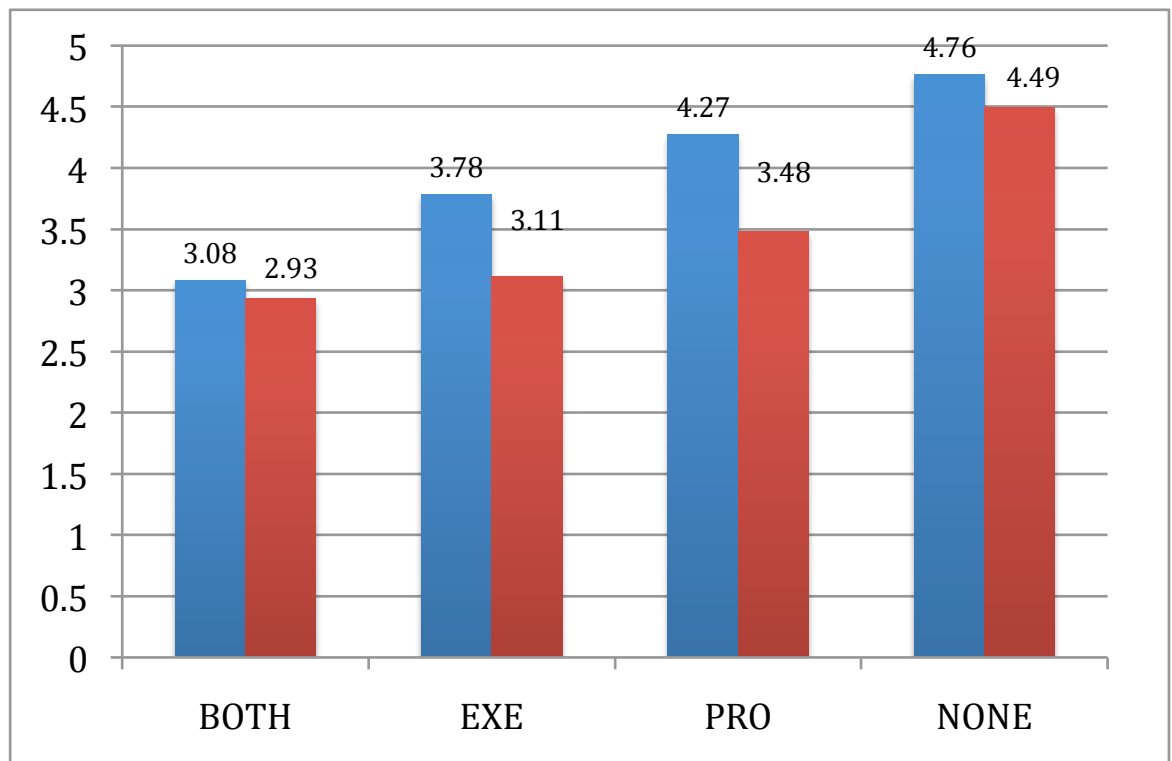
The results of the perception test are in line with the results of the production test. In the production test we have observed the highest long vowel production rates in BOTH set and here in perception test also the highest tendency for long vowels is observed in BOTH set. The order was BOTH (53%) > EXE (31%) > PRO (21%) NONE (0.8 %) in production test, and in perception it is exactly the same order if we compare ratings of the long versions of the words; BOTH (4.56) > EXE (4.17) > PRO (3.77) > NONE (2.63). This shows that the speakers and hearers are influenced by the frequency types in a similar fashion.

Graph 6 and 7 show the relation between the syllable numbers of the items and the ratings of the hearers.



Graph 6. Ratings for Long Versions for Bisyllabics and Trisyllabics

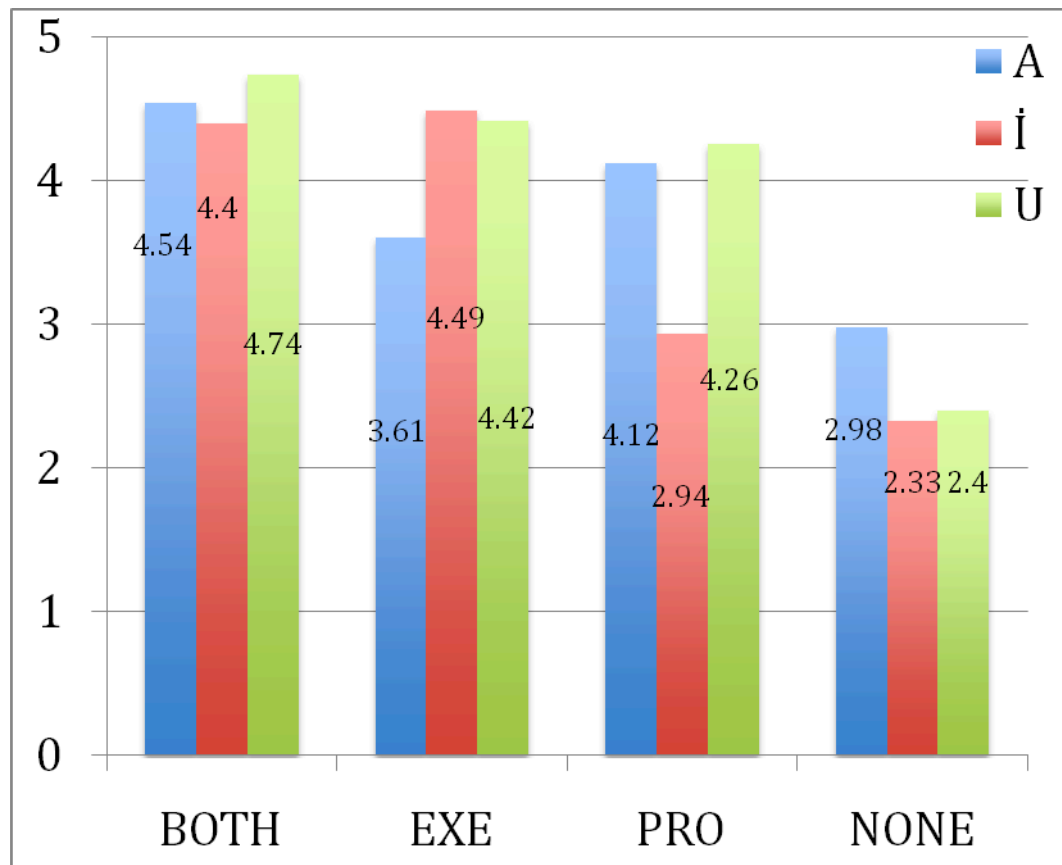
This graph shows a comparison between trisyllabic and bisyllabic nonce items in terms of their well-formedness ratings when they are produced with long vowels. Here we see a consistent rise in ratings in trisyllabics as opposed to bisyllabics. The difference is not high, however it is present in all sets of words consistently. This result is also compatible with the results of the production test. Trisyllabic nonce items were produced and perceived better with long vowels compared to bisyllabics.



Graph 7. Ratings for Short Versions for Bisyllabics and Trisyllabics

Graph 7 is the mirror image of the Graph 6, since one shows the ratings for long versions and the other for the short versions. As expected, in this case bisyllabics are rated better in each group. This is another evidence for the fact that trisyllabics favor long vowels more than bisyllabics.

If we look closely at the results we see that vowel types in each set behave differently.



Graph 8. Rating for long versions for /a/, /i/, /u/

In EXE set /a/ is not rated as high as /i/ and /u/ when they are produced as long. This is the exact same picture that we get in the production test. The explanation that we gave in production test can be valid for this situation as well. The existence of close phonological neighbors with short vowels of the nonce items with /a/ (DAMAS, *damak*, *dama*, etc.) can explain low ratings of EXE items with long /a/. Another striking difference is in the PRO set, the words with long /i/'s are rated very low compared to /a/ and /u/. This is consistent with the production experiment as well. The shortening of /i/ will also be discussed later in the pronunciation survey.

To sum up, these rates suggest that although all forms conform to the phonotactics of Turkish, there may be differences in their well-formedness, and the judgments are directly influenced by the frequency effects regarding the long vowels in Turkish.

5.4 Variation revisited

Results of the production and perception experiments decidedly demonstrate that different frequency effects, i.e., frequent patterns and specific existing items, play a role in the production and the perception of novel items with long vowels. The next question to ask is whether speakers and hearers of Turkish are influenced by these frequency effects when they produce and perceive existing items with long vowels. In 2.4 it was mentioned that in many words variation in vowel length is observed among speakers. The variation can be seen again below. In order to investigate the nature of this variation a pronunciation survey was conducted with two different age groups. First the variations and the patterns in the alternating forms are described.

5.4.1 Variation Among Speakers

- i. Free Variation
- ii. Unusual Lengthening
- iii. Unusual Shortening
- iv. Variation/confusion in rare/novel items

- i. Free Variation

	A	B
(25)	hayır	ha:yır ‘no’

yarın ya:rın ‘tomorrow’

ii. Unusual Lengthening

	A.		B	
	<i>standard</i>		<i>lengthened</i>	
(26)	marul	[marul]	>	[ma:rul] ‘lettuce’
	nasip	[nasip]	>	[na:sip] ‘portion’
	bayan	[bayan]	>	[ba:yan] ‘mrs.’
	hakem	[hakem]	>	[ha:kem] ‘referee’
	tuvalet	[tuvalet]	>	[tuva:let] ‘toilet’
	akraba	[akraba:]	>	[akra:ba:] ‘relative’
	alfabe	[alfabe]	>	[alfa:be] ‘alphabet’
	demokrasi	[demokrasi]	>	[demokra:si] ‘democracy’

iii. *Unusual*¹⁴ Shortening/ Weakening of Vowel Length

	A		B	
	<i>standard</i>		<i>shortened</i>	
(27)	akide	[aki:de]	>	[akide] ‘a type of candy’
	elbise	[elbi:se]	>	[elbise] ‘dress’
	telif	[te:lif]	>	[telif] ‘copyright’
	defile	[defi:le]	>	[defile] ‘fashion show’
	endişe	[endi:şe]	>	[endişe] ‘worry’
	gariban	[gari:ban]	>	[gariban] ‘poor’
	aşiret	[aşı:ret]	>	[aşiret] ‘tribe’

¹⁴ They are unusual from a prescriptive account.

iv. Variation/confusion in rare/novel items

(28)	nakipoğlu	[nakipo:lu]	[na:kipo:lu]
	ergani	[ergani]	[erga:ni]
	daren	[daren]	[da:ren]
	vaniköy	[vaniköy]	[va:niköy]

In her discussion on Tagalog, Zuraw (2000) shows/ argues that exceptions in a language display patterns to make them easier to learn. If we consider the words with long vowels as constituting exceptional set due to their special characteristics such as all being loanwords, etc. we would expect to see some patterns in this set. These patterns were given in Chapter 4. Now let's see whether the variation can be explained with the help of these patterns.

In the data presented above some patterns have emerged. In what follows we will discuss what the emerging patterns suggest about factors that may influence the representation of long vowels in Turkish.

i. In the set where we observe unusual lengthening, in all of the words, the vowel that is subject to lengthening turns out to be /a/. However, in the examples where the process is “unusual shortening”, except for one case of /e/, the vowel that is subject to change is /i/. Actually we have seen in 4.2.1 that /a:/’s outnumber /i/’s. Also in the production and perception tests in PRO set we have seen participants favor short /i/’s. These results add to the idea that long /i/’s display a tendency to get shortened in Turkish.

iv. In words with unusual lengthening, vowels following the lengthened vowel also display a pattern. Despite frontness/backness harmony of vowels in Turkish,¹⁵ the vowels following /a:/ are mostly front vowels. The results in 4.2.1 display the same pattern in the words with long /a/'s. This suggests the relationship between variation and pattern frequency.

5.5 Pronunciation Survey

Besides these two experiments with nonce words, we also had a pronunciation survey with some of the existing words with long vowels in Turkish. The aim of this survey has been to investigate the nature of variation in existing items and see whether any pattern reveals itself regarding the variation observed and whether frequency effects can explain these patterns. One aspect of the variation that we are after is diachronic variation. From the results of the production and perception experiments and the variation data we speculated that /i/'s do not favor length and they may have a tendency to get shortened through time. In order to understand that nature of this kind of diachronic change we had two age groups in this survey.

5.5.1 Participants

Two age groups of 40 monolingual Turkish speakers participated in this survey. The mean age for group 1 (G1) is 78 (n: 20) and for group 2 (G2) is 20 (n: 20).

¹⁵ Turkish displays both internal and external vowel harmony (i.e across affixes). Front vowels /i, e, ö, ü/ are followed by front vowels and back vowels followed by back vowels.

5.5.2 Materials

There were three sets of words in this survey. We went through the list with long vowels and chose the words that we expect that will show variation. The items are chosen with the help of native speaker judgments.

- i) Words that are expected to have unusual shortening in the long vowel (22 items): *elbise, akide, defile* etc.
- ii) Words that are expected to have unusual lengthening (9 items): *marul, demokrasi, tuvalet, alfabe* etc.
- iii) Words that are expected not to change although they have a long vowel (27 items): *ilan, mülakat, nadir, suret* etc.

5.5.3 Procedure

The participants were given the list of words in a random order and asked to read out the words.

5.5.4 Results

Variation in different levels is observed in various items and between two age groups. Each word behaves in a specific manner therefore we did not collapse them into sets and report the variation cumulatively. For example the word *endişe* is produced with a short vowel by 31 out of 40 of the participants , while /i/ in *dakika* is shortened only twice. We analyze each word individually.

The items that show the most drastic variation are shown below in Table 14. First number is the number of participants who produce the forms deviant from the

standard and the numbers in the parentheses represent the percentage rate of the deviant forms over the total in each group.

Table 14. Words that are Subject to Unusual Shortening

	<i>G1</i>	<i>G2</i>	<i>G1+G2</i>	<i>G2-G1</i>
TEYİT	7 (%35)	20 (%100)	27	13
RENCİDE	10 (%50)	20 (%100)	30	10
HEMŞİRE	5 (%25)	15 (%75)	20	10
NAÇAR	5 (%25)	15 (%75)	20	10
ENDİŞE	11 (%55)	20 (%100)	31	9
ELBİSE	10 (%50)	19 (%95)	29	9
BAHARAT	10 (%50)	18 (%90)	28	8
RAKIM	5 (%25)	13 (%65)	18	8
TELİF	8 (%40)	13 (%65)	21	5
AVİZE	7 (%35)	12 (%60)	19	5
GARİBAN	12 (%60)	16 (%80)	28	4

The words *endişe*, *rencide*, *elbise* and, *gariban* are the ones with the highest rate of deviance, i.e. they are unusually shortened hence, can be considered as subjects of a diachronic change. We think the different behavior of two age groups appear to provide evidence for this change. While G1 uses the words with a long vowel, in other words, versions that are given in the dictionary, G2 produces these words with short vowels almost always. This behavior makes us think that these words are in the process of changing/ have been undergoing a change (are subject to change). In some

words the change seems to be more complete than others. For example, *teyit* also seems to be undergoing a change however, it is still produced mostly with long vowels in G1, which suggest the process of change is not as complete as in the word *gariban* or *endişe*.

An interesting result about the change is the fact that we can observe some patterns in the vowels that are shortening. Except for the vowels in the words *naçar*, *rakım* and *baharat* all the vowels shortened are [i]'s and [e]'s. If we look at these three words, *naçar*, *rakım* and *baharat*, we see in some aspects there is a deviation from the prototypes. Especially the vowel sequences are non-prototypical for long vowels, being [a]-[a] and [a]-[ɪ]. [i]'s that are subject to shortening are also followed by [e]'s (a harmonic vowel) which are not prototypical according to the distribution. Another vowel that seems to be shortening in two examples is [e] in *telif* and *teyit*. In these words also we see harmonic vowels [e]-[ɪ]. However [e] in a very similar word *temin* is not shortened by 40 participants. Only one person first uttered the short version after she heard it she immediately changed it to [te:min] with a long [e]. Since *telif* and *temin* have very similar structure, this different behavior towards these words can not be explained by frequent patterns and should be explored further. The fact that both have close neighbors with short vowels, *emin* and *elif* rules out the neighbourhood effect. Another frequency effect we can consider is the token frequency of these items. Bybee (2003) suggests two contradicting effects of token frequency on phonological and morphological change. First, she argues that frequent words are more sensitive to phonetic change for example reduction, i.e. phonetic changes affect the frequent words with a faster rate (Hooper 1976).

However the second effect she suggests works in an opposite way, she argues that more frequent items are more resistant to change since they are more strengthened in the memory than the infrequent words. Therefore frequent words actually resist the regularization process in the language. Since long vowels form a more restricted set in Turkish compared to words with short vowel versions we can say shortening is a regularization process. When we compare the frequency of *temin* and *telif*, in Göz's (2003) frequency dictionary token frequency of *temin* is 76 while the token frequency of *telif* is only 11.¹⁶ The different behaviors of these very similar items lie in the difference in token frequencies. This suggests not only one frequency effect like frequency of patterns or phonological neighborhood but also token frequency of the item may account for the change and difference in the representations.

Another type of change we have observed regarding vowel length is unusual lengthening. The items that have lengthened vowels are actually smaller in number. The results for these words can be seen in Table 15.

Table 15. Words that are Subject to Unusual Lengthening

	<i>G1</i>	<i>G2</i>	<i>G1+G2</i>	<i>G1-G2</i>
AYAR	13 (%65)	0	13	13
MARUL	10 (%50)	0	10	10
HAKEM	3 (%15)	0	3	3
DEMOKRASİ	3 (%15)	0	3	3
NASİP	2 (%10)	0	2	2

¹⁶ *Temin*, *temin etmek* is considered for the number 76 and *telif hakkı* is considered for the number 11.

BAYAN	1 (%5)	0	1	1
AKRABA	1 (%5)	1 (%5)	2	0
TUVALET	0	1 (%5)	1	-1
ALFABE	4 (%20)	9 (%45)	13	-5

Only for the words *ayar marul* and *alfabe* we can argue for a change since others show a small variation. Actually these results are puzzling because we see lengthening in G1 not in G2 for the words *marul* and *ayar*. This would suggest that long forms are the earlier forms, however in the dictionary TDK (1974) these two words are not represented with long vowels. *Alfabe* is also an interesting case, since we see lengthening in G2 more than G1. It can be considered as a hyper-correction behavior, since the word looks like non-native to Turkish, G2 tends to lengthen the [a] which is followed by an [e] and which is in the penult position. Although there is not much variation in this set as the shortened set, still we observe some patterns. For example all the vowels that are subject to unusual lengthening are [a]. Also in other examples like *rakip*, *bakiye*, *viyadük* we observe [a]’s are lengthened by some speakers. There is only one case other than [a] that is *börek* where we see variation in [ö]. There are only a few words with long [ö] in Turkish all of which are signaled in the orthography like *öğren-*, *öğret-*, *öğle* and words that are derived from these words. These structures are very similar to the structure of *börek*, they can be considered as close neighbors, and therefore we can speculate that language use and frequency affect the process of variation in *börek*¹⁷

¹⁷ However since we have a very small sample it may not be a very convincing suggestion. This alternation of [börek]-[bö:rek] remains as a question for now.

Thirdly in the survey we encountered some variation that we would not predict. *Avize* [avi:ze] is produced as [a:vize] by five people, *aşiret* [aşı:ret] is produced as [a:şiret] by two people from G1 and *akide* [aki:de] is produced as [a:kide] by four people from G2. Although these were not expected this variation is not surprising if we take into account the frequent patterns. In these examples we have [a]'s lengthened when they are followed by [i]'s, which is a very frequent pattern as we have mentioned before.

CHAPTER 6

DISCUSSION & CONCLUSION

6.1 Introduction

The aim of this study has been to answer the following questions.

- i. Does language use have an effect on the processes of production and perception of long vowels in Turkish?
- ii. Is there a relation/correlation between the distributional patterns and the production/perception of novel items?
- iii. Do exemplars have an influence on the representation of the words with long vowels?
- iv. Do regularities in the distribution of vowel length have any influence/significance in the variation observed among speakers?

Two experiments with novel items and one pronunciation survey with existing words were designed and conducted among two different age groups to investigate the relation between usage and linguistic processing in the production and perception of long vowels and the variation they exhibit in Turkish.

These experiments and the pronunciation survey have shown that linguistic processes are not independent from usage regarding vowel length in Turkish. In the production experiments it was observed that the novel words which share frequent patterns with the existing words with long vowels in Turkish are more likely to be produced with

long vowels than the words that do not share these patterns. (21% vs. 0.8). Not only frequent patterns but phonological neighborhood was also found to be influential in the production of long vowels. The words that are close phonological neighbors of the existing items tend to be produced with long vowels, even if they do not share frequent patterns. (31%). The results demonstrate that these two effects are independent because when we combine these effects in one set of words, we observed their cumulative influence on production (53%).

The same types of frequency information and the same items are used in well-formedness judgments, and the results from this experiment confirmed the results of the first experiment. The novel words are rated better with long vowels when two kinds of frequency effects are present (4.56/5), this is followed by the neighborhood effect (4.17) and prototype effect (3.87).

Finally the pronunciation survey has shown that vowel length in certain words is subject to change. In some words like *elbise* and *gariban*, change seems to progress faster than some others like *teyit* and *avize*. The frequency effects are also observed in the variation. Especially the vowel type and the vowels that are following the long vowels have conformed to the patterns that we have derived from the lexicon. /a/'s followed by an /i/ and /e/ tend to lengthen and /i/'s followed by a /i/ or /e/ are tend to shorten. Hence we can conclude that the variation process is not independent from usage, the lexicon itself.

The results of the experiments and the survey not only stand as evidence for the existence of frequency effects, but they also contribute to the discussion of the concepts of gradience, redundancy and implicit knowledge. In the following parts of this chapter the implications of the results will be discussed. Later the limitations of the study will be pointed out.

6.2 Nature of the Implicit Knowledge

One of the questions that this study attempts to answer has been what do the language users know implicitly about the phonology of their language, more specifically what do the Turkish speakers know implicitly about vowel length in Turkish. The results of the experiments show that speakers and hearers are aware of the frequent patterns in the input and they use these patterns productively. In the case of Turkish vowel length, the present study has shown language users' awareness of patterns of following and preceding consonants and following vowels of the long vowel. For instance, two nonce items from the experiment, KİLANİ and KAVAPA basically differ from each other in terms of vowels and surrounding consonants of the vowel that we expect to be lengthened. If we look at the production rates for these two specific items, KAVAPA is never produced with a long /a/, however, KİLANİ is produced with a long /a:/ 22 times in 40 utterances. This fact shows the awareness of speakers of the frequent patterns regarding vowel length. To be clear, we should note that LANİ or ANİ sequence in Turkish does not force a long /a/ in absolute terms. KİLANİ with a short /a/ is also perfectly good in Turkish. As seen in (29) Turkish has words with similar sequences with a short vowel.

- | | | | |
|------|---------|-----------|--------------|
| (29) | melanin | [melanin] | ‘melanin’ |
| | mekanik | [mekanik] | ‘mechanical’ |
| | organik | [organik] | ‘organic’ |

These examples show that there is no phonetic or phonotactic constraint in Turkish that will force KİLANİ to be produced as [kila:ni]. However still more than half of the speakers produce it as such. The same logic applies to the opposite example. VAPA sequence does not necessitate a short /a/ in absolute terms. There are not identical but similar words with long /a/ in Turkish.

- | | | | |
|------|---------|------------|--------|
| (30) | bedava | [beda:va] | ‘free’ |
| | harabat | [hara:bat] | ‘ruin’ |

The behavior of speakers supports that the speakers possess a kind of ‘knowledge of tendencies’ in Turkish, that they do not only have discrete rules in mind that will say what is possible and what is impossible but they also have more probabilistic information that would say although both are possible words in Turkish one not the other sounds better. Actually as Bod et al. (2003) argues we can unify this account saying that all the rules are probabilistic; since impossibility and certainty are also probabilities (0 and 1 respectively). In some cases in which a specific rule reveals itself we can say that some tendencies are very strong so they would converge to 0 or 1 and give us the rules that we would consider as absolute.

Apart from the frequent patterns, another implicit knowledge the speakers seem to be making use of productively is the knowledge of specific words in the lexicon. The experiments in this study have shown that Turkish speakers are aware

of the possible lexical neighbors of the nonce words implemented and they are influenced by this information when producing and perceiving novel items. In a sense there is an associative network in the brain and nonce words appear to have activated possible links. This is in line with the Bybee's (2003) model of the lexicon where phonological neighbors are mapped together as a part of “associative network.”¹⁸ This influence of exemplars in Turkish vowel length production and perception also provides evidence for the ‘analogical models’ of language. Many other phenomena in languages like Dutch, German and Spanish are studied under this view (Krott et. al 2001, 2007; Eddington 2000, 2001) and we believe the present study also contributes to the literature on analogy, because it shows that an unpredictable phenomenon can be predicted to some extent on the basis of analogy.

What mechanisms the speakers make use of to acquire this probabilistic information, is an intriguing question. This study proposes frequency as a mechanism and investigates two types of frequency; one being frequent patterns derived from the lexicon, other is the lexical item as a whole unit instead of looking at the parts and patterns. Since patterns are derived from specific words these two effects seem to be interrelated and in fact they are. Language users, of course, do not need to separate these types of information and probably use these effects simultaneously, however, when we separate these types of information for the sake of research we have observed that they are independently effective. When we compare these two effects, according to the results we can conclude that both effects are present and significant in linguistic processes like perception and production, however, the exemplar effect is more prominent than more general patterns.

¹⁸ Bybee (2003) not only suggests a network based on phonological neighborhood but also a network based on semantic relatedness. However in this study we did not investigate the role of semantics. By using nonce items we tried to avoid any effect of semantic relatedness.

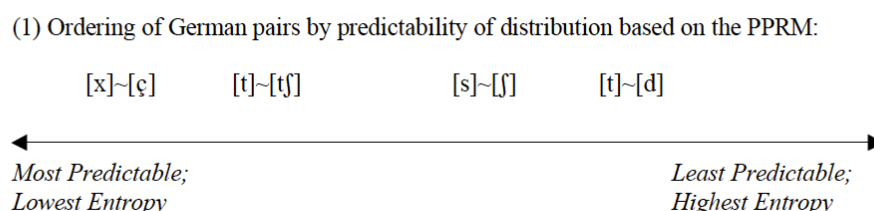
6.3 Gradience

This study being a first attempt to understand the frequency effects in linguistic processes touches upon the issue of gradience in linguistic items and processes as well. Gradience of categories and judgments has been widely studied in recent years. We believe this study will also add to these studies of gradience. First of all, in the present study we showed that tendencies in the lexicon effect the production and perception of long vowels in Turkish. Tendencies are gradient by nature and their influence on processes like perception and production in the case of Turkish vowel length confirms that phenomena like statistical distributional properties can have an effect on linguistic processes. Gradience is not only appreciated/studied by researchers who work on probabilistic linguistics but also some frameworks of Optimality Theory try to reconcile gradience in categories and behavior of language users and language learners. (Boersma 1998; Boersma&Hayes 2001) Variation and optionality in languages and gradient nature of the well-formedness judgments led researchers to introduce a model that use probabilities, Stochastic Optimality Theory, where gradience is sustained with the probabilities assigned to the constraints. This study is not intended to be a study in Stochastic Optimality Theory, however the data on variation (2.4) and the results of the empirical study can be studied further within this framework, as well.

The results of this study can also supplement the discussion of the nature of phonological categories. As stated in 2.2, the vowel length phenomenon that we have studied is defined as lexically specified/unpredictable type of vowel length in Turkish. However the well-formedness judgments and the results of the production

task in this study demonstrate that if a specific environment is provided we can trigger vowel length with nonce words, in other words, we can predict vowel length to some extent. This contradicts with the idea of unpredictability of vowel length, which is a problem for classical categories like phonemes. This behavior of vowel length suggest that to understand the nature of the category of the long vowels in Turkish we cannot rely on only two discrete concepts like phoneme and allophone. In order to solve this kind of problems in languages, Hall (2009) proposes a Probabilistic Phonological Relationship Model (PPRM) in which she takes allophony and contrast as two end points of a continuous scale and tries to quantify the predictability of two sounds on this scale. She builds a highly mathematical model using distributional information, type and token frequency of the distribution of two sounds. These different types of information are combined and quantified using entropy calculations. At the end, the relationship between two sounds is determined to be somewhere on the allophony-contrast scale. For example Hall (2009) investigates consonants in German and the analysis of type and token frequencies reveals results that can be seen in Table 16.

Table 16. Predictability of German Consonant Pairs (Hall, 2009)



This quantification of short-long vowel distinction in Turkish in PPRM is beyond the limits of the current study. However it can be implemented as the following step to

understand the status of vowel length in Turkish since the empirical data that we have laid out suggests that the vowel length in Turkish is not totally unpredictable or predictable. The long vowel-short vowel distinction will be on the continuous scale of phonological relationships (probably closer to the contrast).

6.4 Redundancy

Another difference between classical theories and usage-based model(s) is in the redundancy in representations. Usage-based models favor redundancy in the representation, for example the same word may have two different mental representations and one of them may win the competition (sometimes both are equally good). The behavior of the speakers both in production and perception tests shows that this is the case for vowel length in Turkish. Especially in the perception test in some words both forms have been considered almost equally good, there is no clear preference for one form. Another form of redundancy that we can mention is the redundancy in generalizations that we have derived from the patterns. As stated before, usage-based models lack the idea of abstract economical rules that are independent of linguistic experience. However, we have shown that we can derive generalizations about long vowels in Turkish to some extent, such as; the long vowels are mostly /a:/s, the words with long vowels usually do not conform to the frontness-backness harmony, the surrounding consonants of the long vowels are usually sonorants, the long vowel is usually situated in the penult. In the experiments we have shown that speakers and hearers are implicitly aware of these tendencies and they use them productively. However we have not checked the importance and influence of each tendency. We have a production rate of 21 % in PRO items for

example, and this is the combined effect of these tendencies. The contribution of each tendency may be different and some may be really small, nonetheless we see this influence when they are combined. This combined effect of different properties of prototypical words is an example of redundancy. We cannot conclude from this study that the vowels that are followed by sonorants but that do not employ other frequent patterns that we derived from words with long vowels, are likely to be produced or perceived with long vowels, because we have not checked/examined the influence of the patterns independently, that is why for now we can conclude that we get lengthening when all the properties combined. How much each property adds to the prototype and whether they have a lengthening effect individually or combined can be investigated further with other experiments.

6.5 Token Frequency

Until now the effect of frequency of patterns to the variation and language change is discussed in this survey. We have used type frequency to determine the frequency effects, both in experiments and the discussion.

However token frequency is also discussed in the literature. We have also checked token frequency of the words that we used in the survey to see whether there is a pattern or not. The token frequency results are obtained from Göz (2003)

Table 17. Token Frequency for Unusual Shortening

	G1	G2	G1+G2	G2-G1	Token Frequency
TEYİT	7	20	27	13	7
RENCİDE	10	20	30	10	3
HEMŞİRE	5	15	20	10	33
NAÇAR	5	15	20	10	0

ENDİŞE	11	20	31	9	73
ELBİSE	10	19	29	9	128
BAHARAT	10	18	28	8	30
RAKIM	5	13	18	8	1
TAHİN	2	8	10	6	3
TELİF	8	13	21	5	17
AVİZE	7	12	19	5	13
GARİBAN	12	16	28	4	11
AŞURE	1	5	6	4	2
AKİDE	3	5	8	2	3
AŞİRET	2	4	6	2	10
DAKİKA	0	2	2	2	426
HALİFE	9	10	19	1	22
DEFİLE	4	5	9	1	21
HAZİRAN	2	3	5	1	73
HAZİNE	0	1	1	1	22
NETİCE	3	3	6	0	69
HEZİMET	1	1	2	0	1
ECZANE	0	0	0	0	0
DEFİNE	5	4	9	-1	7
MESİRE	4	3	7	-1	4
HAKİKAT	2	1	3	-1	37
MİRAS	3	0	3	-3	69

Figures show that there is no correlation between token frequency and the variation.

For example the words *elbise* and *rencide* behave very similar in terms of shortening among different age groups, however their token frequencies differ drastically being 128 and 3 respectively. There seems to be no correlation between variation and token frequencies in this case if we compare these words. However, this is a very preliminary result, another survey can be designed to test the effect of token frequency since it was not the aim of this study to measure the token frequency effects.

Table 18. Token Frequencies for Unusual Lengthening

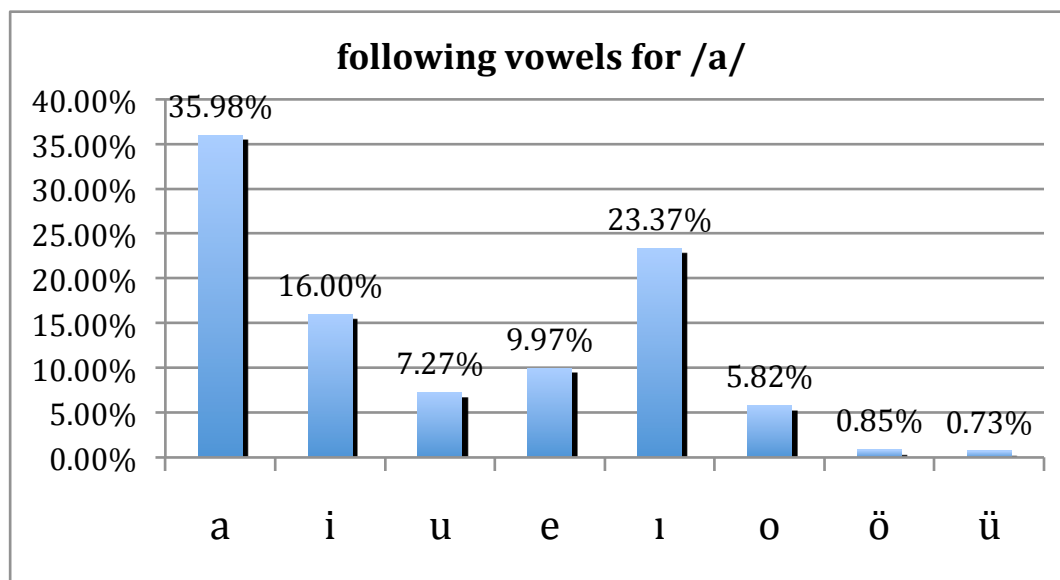
	G1	G2	G1+G2	G2-G1	Token Frequency
ALFABE	4	9	13	5	22
TUVALET	0	1	1	1	59
AKRABA	1	1	2	0	73
BAYAN	1	0	1	-1	180
NASİP	2	0	2	-2	16
DEMOKRASİ	3	0	3	-3	189
HAKEM	3	0	3	-3	63
MARUL	10	0	10	-10	24
AYAR	13	0	13	-13	153

The token frequencies for the unusual lengthening examples also do not suggest a correlation. In order to reach a firm conclusion we can make a list of frequent and infrequent items and repeat the task with more participants. For now these results only display some patterns in the words that are subject to change.

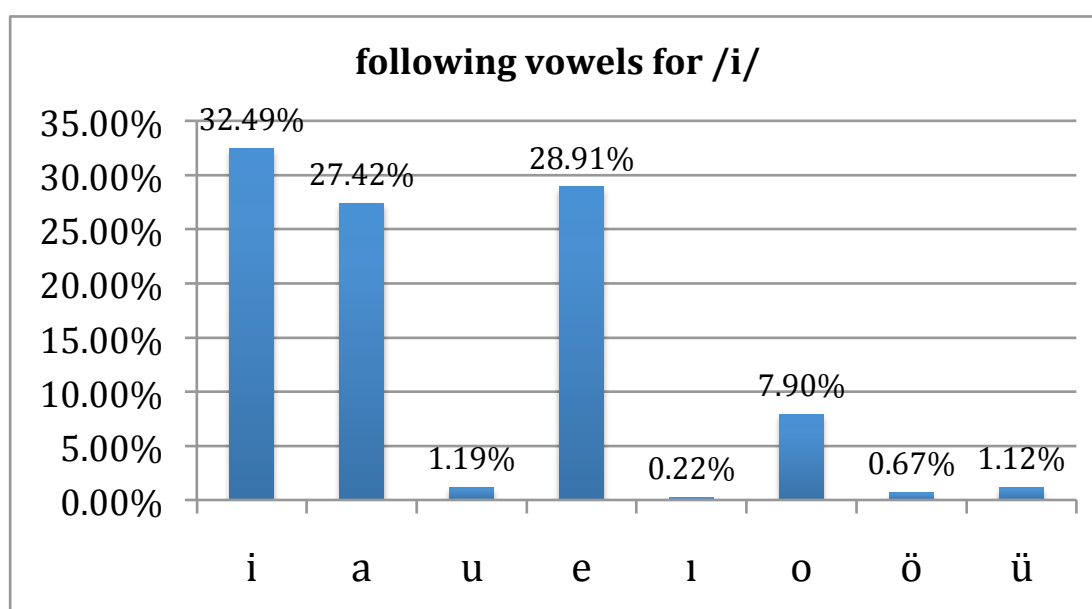
We have investigated the frequent patterns in words with long vowels in Turkish and we have seen that these patterns affect linguistic processes. However we have only analyzed the words with long vowels. A similar analysis of words with short vowels and comparison of these patterns will also help us to have a clearer understanding of most influential patterns. For example we have stated that long vowels are mostly surrounded by sonorants, if this is also true for short vowels than this property alone cannot determine the likelihood of lengthening. We could not carry out such analysis because of time constraints. However we can say that at least one of the properties that we have derived is unique to words with long vowels rather than short vowels: the following vowel. In Turkish in many roots we have a back vowel followed by a back vowel and a front vowel followed by a front vowel, although we do not have an extensive analysis for all the words in Turkish we can

report the results of an analysis of 10259 bisyllabic words compiled from TDK.

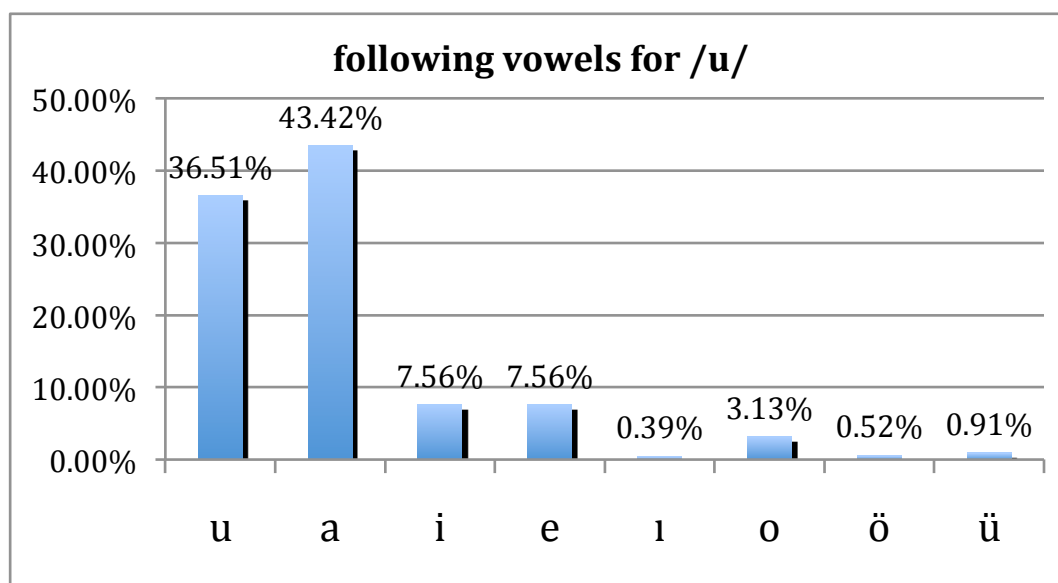
(Nakipoğlu&Kaya, in preperation)



Graph 9. Distribution of Following Vowels for /a/



Graph 10. Distribution of Following Vowels for /i/



Graph 11. Distribution of Following Vowels for /u/

These figures show that /a/ is followed by an /a/ or /ı/ most of the time, /i/ is followed by an /i/ and /e/ and /u/ is followed by an /u/ and /a/. However we had the opposite results for the words with long vowels; /a:/ followed by an /i/, /i:/ followed by an /a/ and /u:/ is followed by an /i/. Therefore we can conclude that the following vowels of the prototypical words with long vowels are unique to words with long vowels, not a general property of Turkish. This contrast seems to be part of the knowledge of Turkish speakers. An extensive analysis with all the words in Turkish and all the properties that we discussed should be done to have a more complete picture.

6.6 The Relationship Between Phonetics and Frequency

Last point we want to discuss is the relation between phonetics and frequency.

Whether the frequent patterns are a direct result of phonetics is a valid question. We derived some main patterns, one being the following vowels and one being the surrounding consonants. Phonetics may be involved in the following consonant

analysis, that is the behavior of the speakers may not be a result of frequency of patterns but some universal phonetic tendencies. For English, for example it has been shown that when a vowel is followed by a voiced fricative it is longer than when it is followed by voiced stops, nasal stops, voiceless fricatives and voiceless stops (House and Fairbanks, 1953). The frequency results of our analysis are not consistent with this order, we had nasal stops and liquids as the most frequent consonants. Another point is that in the data we have analyzed, in almost none of the words the long vowel was in the same syllable with the following consonant unlike the English data. Being in a different syllable may diminish the reported effect of consonants. Therefore the effect of surrounding sonorants seems to be a frequency effect rather than an intrinsic phonetic property in Turkish. There is no phonetic relation of the type of the vowel of the following syllable and the length of the vowel, but it was clearly an important part of the prototype in the present study. Although there might be an influence of phonetics in production and perception of vowel length in Turkish, there are also frequency effects independent from these phonetic tendencies; we cannot say all the frequent patterns that we have derived from the lexicon are directly predicted by phonetic properties of long vowels.

6.7 Conclusion

To conclude, this study has been an attempt to find out the frequency effects in the production and perception of long vowels in Turkish. To that aim it has first investigated the nature of vowel length by carefully studying the distribution of long vowels in the Turkish lexicon. Exhaustively scanning the compiled data for long vowels, it was found out that words that have lexically specified long vowels in

Turkish constitute a special set in terms of the distributional regularities they display and speaker variation. Long vowels, in particular [a:], [i:] and [u:] have been observed to favor mostly bisyllabic and/or trisyllabic words, the penult position in the syllable and the neighboring effects of sonorants. These vowels are further observed to give rise to variation among speakers. In order to find out whether these patterns have any role on the processes like production and perception we have conducted two experiments and a survey, which were designed with the insights gained through the analysis of the data. The results of the experiments and survey suggest that vowel length phenomenon in Turkish is influenced directly from the linguistic experience. If we have a group of words with frequent patterns that are observed in the words with long vowels in Turkish, these words are more likely to be produced with long vowels and they are rated better when they are produced with long vowels compared to the words that do not share these frequent patterns. The same behavior is observed with the words with phonological neighborhood effects. These words tend to be produced with long vowels and they are rated better with long vowels compared to short vowels. Finally when these two individual frequency effects are combined we get the highest rates of long vowel production and highest rates of well-formedness when produced with long vowels. We confirmed the independent effects of two types of frequencies resulting from the language use. This study despite its limitations has successfully demonstrated the influence of language use in linguistic processes. In order to fully understand the effect of frequencies, relative frequencies should be analyzed further. This will be possible if both short vowels in Turkish words are analyzed with a probabilistic view. Another point is that the token frequencies of the existing items can be investigated further. We have used

the frequency dictionary of Göz (2003), It is only based on written documents, however, a frequency analysis of the spoken language will contribute to this kind of studies. This study may be expanded with the help of METU Spoken Turkish Corpus Project.

APPENDICES

Appendix A

Nonce Experimental Items

i) PRO (Prototype effect)

L <u>A</u> NİZ	R <u>I</u> LAK	R <u>U</u> NİF
M <u>A</u> RİT	K <u>I</u> MAS	M <u>U</u> RİN
KİL <u>A</u> Nİ	SAR <u>I</u> LA	TUM <u>U</u> Nİ
TEM <u>A</u> Rİ	TİL <u>I</u> MA	KAN <u>U</u> Rİ

ii) EXE (Exemplar effect)

FAT <u>A</u>	hata	[hata:]	‘mistake’
SEC <u>I</u>	feci	[feci:]	‘bad’
DUF <u>A</u> N	tufan	[tu:fan]‘	flood’
D <u>A</u> MAS	damat	[da:mat]	‘groom’
N <u>I</u> BE	hibe	[hi:be]	‘donation’
K <u>U</u> SA	musa	[mu:sa]	‘male name’
F <u>A</u> TIRA	hatıra	[ha:tıra]	‘memory’
<u>İ</u> TİRAT	itiraf/itiraz	[i:tiraf]/[iti:raz]	‘confession/objection’
M <u>U</u> TEBEN	muteber	[mu:teber]	‘authentic’
AKRAD <u>A</u>	akraba	[akraba:]	‘relative’
EZEM <u>I</u>	ezeli	[ezeli:]	‘primordial’
KEBERR <u>U</u>	teberru	[teberru:]	‘bequest’

iii) BOTH (Both exemplar and prototype effect)

N <u>A</u> RÎM	narin	[na:rin]	‘delicate’
N <u>I</u> LAT	milat	[mi:lat]	‘the birth date of Christ’
M <u>U</u> NÎZ	munis	[mu:nis]	‘tame’
Z <u>I</u> MA	sima	[si:ma]	‘face’
D <u>U</u> SE	buse	[bu:se]	‘kiss’
T <u>A</u> RÎN	tarih	[ta:rih]	‘date’
NÎN <u>A</u> RE	minare	[mina:re]	‘minaret’
TAHS <u>I</u> LAK	tahsilat	[tahsi:lat]	‘payments received’
LUM <u>U</u> NE	numune	[numu:ne]	‘sample’
KAHR <u>I</u> BAT	tahribat	[tahri:bat]	‘damage’
PAZ <u>U</u> LET	kazulet	[kazu:let]	‘huge’
CES <u>A</u> REK	cesaret	[cesa:ret]	‘courage’

iv) NONE (Prototypically short)

V <u>A</u> ÇAP	Y <u>I</u> PÎT	Ç <u>U</u> YUK
P <u>A</u> ŞIF	Ş <u>I</u> PEZ	F <u>U</u> ÇAV
TAG <u>A</u> ŞA	KEY <u>I</u> PÎ	TAÇ <u>U</u> YA
KAV <u>A</u> PA	EŞ <u>I</u> YÎF	KUV <u>U</u> PA

Appendix B

The Items in the Pronunciation Survey

i. Items for which we predict unusual shortening

TEYİT	TELİF	HAZİRAN
RENCİDE	AVİZE	HAZİNE
HEMŞİRE	GARİBAN	NETİCE
NAÇAR	AŞURE	HEZİMET
ENDİŞE	AKİDE	ECZANE
ELBİSE	AŞİRET	DEFİNE
BAHARAT	DAKİKA	MESİRE
RAKIM	HALİFE	HAKİKAT
TAHİN	DEFİLE	MİRAS

ii. Items for which we predict unusual lengthening

ALFABE	BAYAN	HAKEM
TUVALET	NASİP	MARUL
AKRABA	DEMOKRASİ	AYAR

iii. Items for which we do not predict any change

TUFAN	HAVADİS	İCAT
BAZEN	MÜLAKAT	İMAN
MASUM	HATA	İLAN
HATIRA	NADİR	DESİSE
ADETA	ADALET	TESİR
AŞİNA	ADİL	SURET
BEDAVA	AYİN	TEMİN
BİRADER	DAMAT	FUZULİ
LAZIM	İTİRAF	HUSUMET

Appendix C

Results of the Distributional Analysis

i. Position of the long vowels

Table 1C. Monosyllabics

STRUCTURE OF MONOSYLLABICS	CVC
12	12

Table 2C. Words with four syllables

STRUCTURE OF WORDS WITH 4 SYLL.	n	%
ANTEPENULT	112	53.59
PENULT	51	24.40
BOTH 1. & 3.	14	6.70
1. SYLL	13	6.22
FINAL	10	4.78
BOTH 2. & 3.	3	1.44
BOTH 2. & 4.	3	1.44
BOTH 1. & 4.	2	0.96
BOTH 1. & 2. & 3.	1	0.48
BOTH 1. & 2.	0	0.00
BOTH 3. & 4.	0	0.00
TOTAL	209	100.00

Table 3C. Words with five syllables

STRUCTURE OF WORDS WITH 5 SYL.	n	%
PENULT	7	43.75
ANTEPEN.	3	18.75
1ST SYL.	2	12.50
2ND SYL.	2	12.50
BOTH 1. & 4.	2	12.50
FINAL	0	0.00
TOTAL	16	100.00

Table 4C. Words with six syllables

STRUCTURE OF WORDS WITH 6 SYL.	n	%
PENULT	2	66.67
FINAL	1	33.33
TOTAL	3	100.00

ii. Following Consonants

Table 5C. Following consonants

A	<i>n</i>	%
a:n	137	11.9
a:r	133	11.5
a:l	110	9.5
a:h	94	8.1
a:b	80	6.9
a:d	80	6.9

Ī	<i>n</i>	%
i:m	28	12.7
i:r	28	12.7
i:l	25	11.3
i:k	18	8.1
i:d	17	7.7
i:t	17	7.7

U	<i>n</i>	%
u:r	26	19.1
u:n	16	11.8
u:d	14	10.3
u:l	13	9.6
u:b	13	9.6
u:t	13	9.6

a:m	70	6.1
a:k	69	6.0
a:y	62	5.4
a:s	61	5.3
a:z	58	5.0
a:t	53	4.6
a:f	47	4.1
a:v	47	4.1
a:c	23	2.0
a:ş	13	1.1
a:g	7	0.6
a:p	6	0.5
a:ç	4	0.3
a:j	0	0.0
TOT.	1154	100

i:n	15	6.8
i:b	12	5.4
i:z	11	5.0
i:c	10	4.5
i:s	9	4.1
i:f	8	3.6
i:ş	6	2.7
i:v	6	2.7
i:h	4	1.8
i:g	2	0.9
i:p	2	0.9
i:y	2	0.9
i:ç	1	0.5
i:j	0	0.0
TOT.	221	100

u:m	7	5.1
u:s	7	5.1
u:k	6	4.4
u:z	6	4.4
u:h	5	3.7
u:f	4	2.9
u:c	3	2.2
u:ş	2	1.5
u:ç	1	0.7
u:g	0	0.0
u:j	0	0.0
u:p	0	0.0
u:v	0	0.0
u:y	0	0.0
TOT.	136	100

iii. Preceding Consonants

Table 6C. Preceding Consonants

A	<i>n</i>	%
ma:	115	9.6
la:	99	8.3
ra:	96	8.0
ha:	93	7.8
ka:	91	7.6
sa:	78	6.5
na:	76	6.4
ta:	74	6.2
ya:	72	6.0
va:	65	5.4
ba:	60	5.0
za:	60	5.0

U	<i>n</i>	%
ru:	20	13.8
mu:	19	13.1
su:	19	13.1
hu:	14	9.7
tu:	12	8.3
bu:	11	7.6
ku:	8	5.5
cu:	7	4.8
lu:	6	4.1
nu:	6	4.1
şu:	6	4.1
gu:	4	2.8

İ	<i>n</i>	%
ri:	49	12.9
ki:	37	9.8
li:	36	9.5
si:	33	8.7
bi:	27	7.1
ni:	26	6.9
di:	25	6.6
zi:	25	6.6
mi:	21	5.5
ti:	21	5.5
vi:	21	5.5
fi:	20	5.3

da:	56	4.7
fa:	40	3.4
şa:	38	3.2
ca:	37	3.1
ga:	23	1.9
pa:	16	1.3
ça:	4	0.3
ja:	1	0.1
TOT.	1194	100

yu:	4	2.8
zu:	4	2.8
du:	3	2.1
fu:	2	1.4
çu:	0	0.0
ju:	0	0.0
pu:	0	0.0
vu:	0	0.0
TOT.	145	100

hi:	11	2.9
şı:	11	2.9
ci:	8	2.1
yi:	6	1.6
pi:	2	0.5
çi:	0	0.0
gi:	0	0.0
ji:	0	0.0
TOT.	379	100

REFERENCES

- Bailey, T. M. and Hahn, U. 2001. Determinants of wordlikeness: Phonotactics or lexical neighborhoods? *Journal of Memory and Language*, 4, 568 – 591.
- Bod, R., Hay, J. and Jannedy, S. eds., 2003. *Probabilistic Linguistics*. Cambridge, MA: MIT Press.
- Boersma, P. 1998. *Functional Phonology*. Ph.D. thesis, University of Amsterdam.
- Boersma, P. and Hayes, B., 2001. Empirical tests of the gradual learning algorithm. *Linguistic Inquiry*, 32(1):45–86.
- Bybee, J., 1985. *Morphology: a study of the relation between meaning and form*. Philadelphia: Benjamins.
- Bybee, J. 2001. *Phonology and language use*, Cambridge U.K: Cambridge University Press.
- Bybee, J., 2006. From Usage to Grammar: The Mind's Response to Repetition *Language* – Vol. 82, No. 4, pp. 711-733
- Bybee, J. and J. McClelland. 2005. Alternatives to the combinatorial paradigm of linguistic theory based on domain general principles of human cognition. *The Linguistic Review*. 22, pp. 381-410.
- Çakır, M. C. 2000. On Non-Final Stress in Turkish Simplex Words. *Proceedings of the Ninth International Conference on Turkish Linguistics*. Vol. 9. No. 9 pp.1-8.
- Dell, G.S., Reed K.D., Adams D. R., Meyer, A.S. 2000. Speech Errors, Phonotactic Constraints, and Implicit Learning: A Study of the Role of Experience in Language Production *Journal of Experimental Psychology: Learning, Memory, and Cognition* 2000, Vol. 26, No. 6, 1355-1367
- Denwood, A., 2002. k - ø : morpho-phonology in Turkish. *SOAS Working Papers in Linguistics and Phonetics* 12, 89-98.
- Eddington, D., 1996 Diphthongization in Spanish Derivational Morphology: An Empirical Investigation, *Hispanic Linguistics*. 8.1-35
- Eddington, D., 2000. Spanish stress assignment within the Analogical Modeling of Language, *Language*. 76. 92-109.
- Eddington, D., 2001. A Usage-based Simulation of Spanish S-weakening, *Journal of Italian Linguistics*. 13. 191-209.
- Ergenç, İ., 1995. *Konuşma Dili ve Türkçe'nin Söyleyiş Sözlüğü*. Simurg, Ankara

- Goldrick, M., and Larson, M., 2008. Phonotactic probability influences speech production. *Cognition*, 107, 1155-1164.
- Göksel, A. and Kerslake, C., 2005. *Turkish, A Comprehensive Grammar*. London: Routledge.
- Göz, İ., 2002. *Yazılı Türkçenin Kelime Sıklığı Sözlüğü* TDK Yayınları
- Hall, K. C., 2009. *A probabilistic model of phonological relationships from contrast to allophony*. Ph.D. Dissertation OSU.
- Hooper, J. Bybee. 1976. Word Frequency in Lexical Diffusion and the source of morphophonological change. *Current Progress in Historical Linguistics*, ed., W. Cristie, 96-105. Amsterdam: North Holland
- House, A.S. and Fairbanks, G., 1953. The influence of consonantal environment upon the secondary acoustical characteristics of vowels. *Journal of the Acoustical Society of America* 25: 105-113
- Inkelas, S. and Orgun, O., 1995. Level Ordering and Economy in the Lexical Phonology of Turkish. *Language* 71.763-793
- Jusczyk, P. W., Luce, P. A., and Charles-Luce, J., 1994. Infants' sensitivity to phonotactic patterns in the native language. *Journal of Memory and Language*. 33, 630-645.
- Kornfilt, J., 1985. Stem-Penultimate Empty Cs, Compensatory Lengthening and Vowel Epenthesis in Turkish, in Wetzels, L and Sezer, E (eds) *Studies in Compesatory Lengthening* (pp.79-96) Foris Publications, Holland.
- Krott, A., Schreuder, R. and Baayen, R.H., Dressler, W.U. 2007. Analogical effects on linking elements in German compounds. *Language and Cognitive Processes*. 22 (1), 25-57.
- Krott, A., Baayen, R. H. and Schreuder, R., 2001. Analogy in morphology: modeling the choice of linking morphemes in Dutch. *Linguistics*, 39 (1), 51-93.
- Langacker, R.W., 1991. *Concept, Image, and Symbol: The Cognitive Basis of Grammar*. Berlin & New York: Mouton de Gruyter.
- Lees, R. B. 1961. *Phonology of modern standard Turkish*. Bloomington: Indiana University Press.
- Lehiste, I. 1970. *Suprasegmentals*. Cambridge, MA: MIT Press.

- Maye, J., & Gerken, L. 2000. Learning phonemes without minimal pairs. *Proceedings of the 24th Annual Boston University Conference on Language Development*, 522 – 533
- Maye, J., Werker, J. F., and Gerken, L. 2002. Infant sensitivity to distributional information can affect phonetic discrimination. *Cognition*. 82 (3), B101–B111.
- Nakipoğlu, M & Kaya, M., in preparation. The distributional analysis of Turkish lexicon
- Nuhbalaoğlu, D. 2010 *On the role of empty onsets in Turkish: a Government Phonology approach*. MA thesis Boğaziçi University.
- Pierrehumbert, J. 2001. Stochastic phonology. *GLOT* 5:6, 1-13.
- Pierrehumbert, J. 2003a. Phonetic diversity, statistical learning, and acquisition of phonology, *Language and Speech*, 46 (2-3), 115-154.
- Pierrehumbert, J. 2003b. Probabilistic theories of phonology. In R. Bod, J. B. Hay, and S. Jannedy eds., *Probability theory in linguistics* (pp. 177 – 228). Cambridge, MA: MIT Press
- Özsoy, S. 2004. *Türkçe'nin Yapısı-I Sesbilimi*, B.Ü. Dil Merkezi Türkçe'nin Yapısı Dizisi. İstanbul: Boğaziçi Üniversitesi Yayınları.
- Rosch, E. 1973. Natural categories, *Cognitive Psychology* 4: pp.328-50.
- Rosch, E. 1978. *Principles of Categorization*. *Cognition and categorization*, eds. E. Rosch and B. Lloyd, pp.27-47. Hillsdale, NJ:Erlbaum.
- Saffran, J. R., Newport, E. L., and Aslin, R. N. 1996b. Word segmentation: The role of distributional cues. *Journal of Memory and Language* 35, pp.606–621.
- Saffran, J. R., Aslin, R. N., and Newport, E. L. 1996a. Statistical cues in language acquisition: Word segmentation by infants. *Proceedings of the 18th Annual Conference of the Cognitive Science Society (COGSCI-96)*, San Diego, Calif., pp. 376–380.
- Sezer, E. 1981 'On non-final stress in Turkish', *Journal of Turkish Studies* 5: 61-69.
- Sezer, E. 1985. An Autosegmental Analysis of Compensatory Lengthening and Vowel Epenthesis in Turkish, in Wetzels, L and Sezer, E eds. *Studies in Compesatory Lengthening* 227-250.
- Taylan, E. 2010. Lexically specified vowel length in Turkish. Manuscript

TDK 1974. *Türkçe Sözlük*, 6. baskı, TDK yayını, Ankara.

Treiman, R. et al. (2000) English speakers' sensitivity to phonotactic patterns, in Broe & Pierrehumbert. *Papers in Laboratory Phonology V: Acquisition and the Lexicon*. Cambridge, UK: Cambridge University Press. pp 269-282.

Vitevitch, M. S. & Luce, P. A. 1999. Probabilistic phonotactics and neighborhood activation in spoken word recognition. *Journal of Memory and Language*, 40, 374–408.

Vitevitch, M. S., Luce, P. A., Charles-Luce, J., & Kemmerer, D. 1997. Phonotactics and syllable stress: Implications for the processing of spoken nonsense words. *Language and Speech*, 40, 47–62.

Zuraw, K. 2000. *Patterned Exceptions in Phonology*. Ph.D. dissertation, UCLA.

Zuraw, K. 2003. Probability in language change. In Rens Bod, Jennifer Hay, and Stephanie Jannedy, (eds), *Probabilistic Linguistics*. Cambridge, MA: MIT Press. pp. 139-176.