

UNIVERSITY OF CALGARY

The Acquisition of L2 Segmental Contrasts:
English Speakers' Perception and Production of
Czech Palatal Stops

by

Susan Barbara Atkey

A THESIS
SUBMITTED TO THE FACULTY OF GRADUATE STUDIES
IN PARTIAL FULFILMENT OF THE REQUIREMENTS FOR THE
DEGREE OF MASTER OF ARTS

DEPARTMENT OF LINGUISTICS

CALGARY, ALBERTA
DECEMBER 2000

©Susan Barbara Atkey 2001



National Library
of Canada

Acquisitions and
Bibliographic Services

395 Wellington Street
Ottawa ON K1A 0N4
Canada

Bibliothèque nationale
du Canada

Acquisitions et
services bibliographiques

395, rue Wellington
Ottawa ON K1A 0N4
Canada

Your file / Votre référence

Our file / Notre référence

The author has granted a non-exclusive licence allowing the National Library of Canada to reproduce, loan, distribute or sell copies of this thesis in microform, paper or electronic formats.

L'auteur a accordé une licence non exclusive permettant à la Bibliothèque nationale du Canada de reproduire, prêter, distribuer ou vendre des copies de cette thèse sous la forme de microfiche/film, de reproduction sur papier ou sur format électronique.

The author retains ownership of the copyright in this thesis. Neither the thesis nor substantial extracts from it may be printed or otherwise reproduced without the author's permission.

L'auteur conserve la propriété du droit d'auteur qui protège cette thèse. Ni la thèse ni des extraits substantiels de celle-ci ne doivent être imprimés ou autrement reproduits sans son autorisation.

0-612-64902-4

Canada

ABSTRACT

THE ACQUISITION OF L2 SEGMENTAL CONTRASTS: ENGLISH SPEAKERS' PERCEPTION AND PRODUCTION OF CZECH PALATAL STOPS

This thesis examines English speakers' perception and production of a non-native segmental stop contrast in Czech: alveolar /t, d/ versus palatal /c, ɟ/. Brown's (1997, 1998, 2000) model of L2 phoneme acquisition argues that it is the L1 feature inventory (rather than individual segments) which define the boundaries within which novel phonemes are perceived. Specifically, L1 speakers can perceive novel L2 contrasts if that contrast is characterised by a feature present in their L1 grammar: conversely, if a particular feature is lacking in the L1 feature inventory, then perception of the novel phonemic contrast should be precluded. I argue that the contrasting feature between Czech alveolar /t, d/ and palatal /c, ɟ/ is [posterior]. English requires the feature [posterior] to contrast two fricative segments: alveolar /s, z/ versus alveo-palatal /ʃ, ʒ/. English speakers thus have the necessary building block for acquisition of the novel Czech contrast.

ACKNOWLEDGEMENTS

Many people have helped in one way or another during the writing of this thesis, and it is a pleasure to finally acknowledge them. First, I wish especially to thank my advisor, Dr. John Archibald, for his advice and support and especially for his patience. Thanks again for encouraging me to present my thesis research at the New Sounds conference in Amsterdam. John -- it was terrifying but so worth it! I would also like to thank my examining committee members Dr. Robert Murray and Dr. Elena Bratishenko. Through their careful editing and detailed criticisms, this thesis has been greatly improved.

I would like to thank the Department of Linguistics and the Faculty of Graduate Studies for providing scholarships and financial support. Thanks are also due to the members of the Linguistics Department, both faculty and students, for providing the environment which made the study of linguistics enjoyable. I want to especially thank my fellow grad student Jana Carson -- Jana, can you believe we actually made it through!

Many people in my personal life have also been crucial to the completion of this thesis. A huge debt of gratitude is owed to my parents, John and Eva Atkey, for believing in me and supporting me through my seemingly never-ending years of school. Thank you for everything! I also want to thank my sister Jill Atkey as well as Dale Rahim for providing laughs and taking me out for air once in a while. I am grateful to Anne and Bert Herman for their support and encouragement. And I cannot forget my friends, who reminded me that there is a life outside of school.

I especially want to thank Julian Duerksen for supporting and putting up with me -- with such patience and good humor! -- through grad school boot camp.

A final word of thanks goes out to my subjects, without whom, after all, this thesis would not have been possible. A huge thank you and best wishes in your Czech-learning endeavours!

TABLE OF CONTENTS

THE ACQUISITION OF L2 SEGMENTAL CONTRASTS: ENGLISH SPEAKERS' PERCEPTION AND PRODUCTION OF CZECH PALATAL STOPS

Approval Page	ii
Abstract	iii
Acknowledgements	iv
Table of Contents	v
List of Tables	viii
List of Figures	ix

INTRODUCTION	1
--------------------	---

CHAPTER ONE: PHONOLOGICAL THEORY

0.0	Background	5
1.0	Feature Geometry Theory	5
2.0	Underspecification theory	12
2.1	Contrastive Specification	13
2.2	Radical Underspecification	14
2.3	Minimally Contrastive Underspecification	15
3.0	Coronal Underspecification	17
3.1	The Special Status of Coronals	18
3.2	The Underspecification of Coronals	19

CHAPTER TWO: L1 ACQUISITION OF SEGMENTAL PHONOLOGY

0.0	Background	24
1.0	First Language Acquisition	24
2.0	Acquisition of the Feature Geometry	28
2.1	Rice & Avery (1995)	31
2.2	Brown (1997, 2000): L1 phonological system in speech perception ..	40
2.2.1	Infant Speech Perception	41
2.2.2	Phonological development and acoustic discrimination	42

CHAPTER THREE: SECOND LANGUAGE PHONOLOGICAL ACQUISITION

0.0	Background	46
1.0	Second Language Phonological Acquisition: An Overview	47
1.1	Access to UG in Second Language Acquisition	48
2.0	Models of Second Language Feature Acquisition	52
2.1	Characteristics of a good model of second language acquisition	52
2.2	Theoretical Models of Second Language Acquisition	54
2.2.1	Perceptual Assimilation Model (PAM) (Best, 1993, 1994) ...	55
2.2.2	Speech Learning Model (Flege, 1991, 1995)	55
2.2.3	Feature Competition Model (FCM) (Hancin-Bhatt 1994)	56
3.0	Brown's Model of L1 Interference	57
3.1	The Phonological Filter	59
3.2	Implications for L2 Acquisition	61
3.3	The experimental data (Brown, 1998)	63
3.3.1	Experiment 1: Japanese and Chinese speakers' acquisition of English /l/ versus /r/	63
3.3.2	Experiment 2: Japanese speakers' acquisition of English /l, r/, /b, v/ and /f, v/	69
3.4	The Influence of Pronunciation Training on the Perception of Second Language Contrasts (Matthews, 1997)	72

CHAPTER FOUR: SEGMENTAL PHONOLOGY OF CZECH AND ENGLISH

0.0	Background	75
1.0	The Segmental Inventory of Czech	75
1.1	Vowels	76
1.2	Consonants	76
1.2.1	Plosives	77
2.0	Czech Coronals	80
2.1	Alveolars	81
2.2	Palatals	81
2.2.1	Palatal articulations	81
2.2.2	Palatals as Coronal sounds	82
2.3	Czech Feature Hierarchy	83
2.3.1	Anterior versus Posterior	84
3.0	The Phonology of English	86
3.1	English Coronal Fricatives	87
3.2	Feature Representations of English Fricatives	88
3.2.1	Feature Representation of Alveo-palatal /ʃ, ʒ/	90
3.2.2	Feature Representation of Dental /θ, ð/	91

CHAPTER FIVE: THE ACQUISITION OF CZECH PALATAL STOPS

0.0	Background	94
1.0	Experiment One: Forced Choice Phoneme Selection Task	95
1.1	Methodology	97
1.1.1	Subjects	97
1.1.2	Stimuli	98
1.2	Results	101
1.3	Discussion	106
2.0	Experiment Two: Production Tasks	108
2.1	Methodology	108
2.1.1	Subjects	108
2.1.2	Stimuli	109
2.2	Results	109
2.2.1	Free Production	109
2.2.2	Sentence Reading Task	115
2.3	Discussion	118
3.0	Stages of Acquisition: Czech Palatal Stops	121

CHAPTER SIX: AVENUES FOR FUTURE RESEARCH AND CONCLUSIONS

0.0	Background	127
1.0	Avenues for Future Research	127
2.0	Conclusion	128

Appendix One	131
Appendix Two	135
Appendix Three	141
References	145

LIST OF TABLES

Table 3.1	Mean Results of Discrimination Task: English /l/ vs /r/	68
Table 3.2	Mean Results of Picture Selection Task: English /l/ vs /r/	68
Table 3.3	Mean Results of Discrimination Task: English Phonemic Contrasts . .	71
Table 3.4	Mean Results of Picture Selection Task: English Phonemic Contrasts .	71
Table 3.5	Relationship Between English Contrasts and Japanese Inventory	73
Table 4.1	Phonemic Inventory of Czech Consonants	77
Table 4.2	Phonemic Inventory of English Consonants	87
Table 5.1	Percentage of Czech Palatal Stop Tokens Perceived Correctly	102
Table 5.2	English Speakers' Production of Palatal Stops in Free Speech	110

LIST OF FIGURES

Figure 2.1	Modular Framework of Language Acquisition	26
Figure 2.2	Model of Phonological Structure	28
Figure 2.3	Mediation of speech perception by phonological structure	44
Figure 3.1	Mediation of speech perception by phonological structure	60

INTRODUCTION

This thesis combines two broad areas of linguistic inquiry: second language acquisition and generative phonology. The merging of two such diverse areas of linguistic inquiry is an attempt to address what is known generally as “Plato’s problem”, or, in linguistic terms, the “poverty of the stimulus argument” (Lightfoot, 1981): How are humans able to acquire such rich and complex systems of knowledge that do not accurately reflect the relatively limited input they are exposed to? This gap between the rapid acquisition of complex grammatical structures and deficient input led Chomsky to posit the innate mental structure known as Universal Grammar (UG), which both constrains cross-linguistic variation while informing the process of language acquisition. Recent models of generative phonology are thus models of phonological acquisition in that children’s constructions of phonological representations are both constrained and guided by principles of UG. Generative phonology addresses the question by dealing with the mental representations of speech sounds (as opposed to the physical or phonetic implementation), while the area of L2 acquisition gives us language learning on which to test the theories. Thus, this thesis addresses two broad questions: (i), whether or not novel segmental contrasts in the L2 can be acquired, and (ii), how speech segments are mentally represented.

The phonological framework of Feature Geometry theory provides a formal means for representing mental phonemic knowledge by assuming that (i), individual speech segments have internal organization composed of sub-segmental features supplied from a finite set provided by UG, and (ii), features are organised hierarchically reflecting both markedness and phonological dependency. Crucially, no language uses all the possible UG-provided features and no segment requires all the features of that language. Sound segments are thus not the primitives of language per se: rather, they are the sum of hierarchically organised features such that two segments may contrast on the basis of one feature alone.

Recent research (Rice & Avery, 1995) in first language acquisition of segmental (sound system) structure has argued that the acquisition of language-specific sound

contrasts is a step-by-step process of feature elaboration. The precise features comprising a segment are acquired in a hierarchical fashion, following the relationships of dependency and constituency encoded in the structure. Structural elaboration can thus account for both variability and uniformity. Cynthia Brown (1993, 1997, 2000; see also Brown and Matthews; 1993, 1997) takes the claims of segmental acquisition as a process of structural elaboration one step further by integrating the L1 phonological system in both L1 and L2 speech perception and production. As Brown's model is essential to assumptions and experimental research in this thesis, I will take some time to explain it here. Based on findings from infant speech perception showing that a decrease in perceptual ability to discriminate *non-native* sounds corresponds to an increase in phonological ability to discriminate *native* segmental contrasts, Brown proposes that the link between phonological development and acoustic discrimination can be accounted for by the same mechanism. Essentially, Brown proposes that step-by-step elaboration of the hierarchical feature geometry in first language acquisition imposes a template or filter on the perceptual system within which the language-specific phonemic categories are perceived. It is the *detection* of phonemic contrasts in the input language which triggers the elaboration of the language-specific phonological hierarchy. The phonological structure then acts as intermediary between the acoustic signal and the linguistic system by channeling the distinct acoustic signals into phonemic categories, guided by language specific featural makeup. That is, the feature structure is used to funnel distinct phonetic variations into individual phonemic categories so that perception of non-native contrasts gradually declines as the novel sounds are interpreted as phonetic variants of existing categories. The segmental representation characterising a phonemic category will be activated by *intra*-category phonetic variants but not by *inter*-category variants, thereby assisting in speech processing by allowing noise without compromising the recoverability of the underlying representation.

What of the L2 learner who arrives at the language learning task with an existing phonological structure from the native language? Brown argues that in L2 acquisition the intake to the language acquisition device is determined by the phonological structure of the first language. That is, the phonological structure of the L1 acts as a sort of template defining the categories within which the L2 sounds are perceived. However, it is not the

segments themselves that determine perception, but rather the segmental sub-units, or features, of segments that play an essential role in the acquisition of L2. Whether or not a L2 learner can perceive the non-native phonological contrasts is dependent on the feature composition of the L1 phonology. Specifically, if a non-native segment is characterised by a particular feature, then the learner will be able to perceive the non-native contrasts if he or she manipulates the feature elsewhere in the L1 grammar. Brown points out that an L2 learner's experience perceiving L1 phonemic contrasts along an acoustic dimension defined by a particular underlying feature permits him or her to accurately discriminate *any* phonemic contrast differing along that same dimension, despite a lack of acoustic, phonetic or phonemic experience with a *particular* non-native contrast. Perception of a new phonemic contrast is facilitated by the presence of the distinguishing feature elsewhere in the learner's L1 inventory. Conversely, if a particular feature is lacking in the representation of any phonemic contrasts in the L1 feature geometry, then perception of the novel phonemic contrast should be precluded. In this case, the filtering of the acoustic signal which aids in L1 processing can negatively influence the perception of a non-native language, as intra-category variation in the L1 may actually constitute phonemic contrasts (or inter-category variation) in the L2.

Brown presents experimental evidence on Japanese and Mandarin Chinese speakers' acquisition (or non acquisition, in the case of the native Japanese subjects) of the English /l/ versus /r/ contrast in support of her hypothesis that it is the featural rather than the segmental level determining whether or not L2 learners will be able to perceive and produce a novel L2 contrast. As mentioned, this work by Brown provides the impetus for the experimental research conducted in this thesis. However, in order to address the broad questions of L2 segmental representation and acquisition, one needs to speak to the specific. For this thesis, I investigated six native North American English speakers' perception and production of a non-native segmental stop contrast: Czech alveolar /t, d/ versus palatal /c, j/. English speakers learning Czech are faced with a non-native phonemic contrast in that Czech contrasts two coronal stop places of articulation: alveolar /t, d/, which do occur in English, versus palatal /c, j/, which do not. In order to determine the likelihood of English speakers' acquisition of this novel segmental contrast,

I needed to establish the distinguishing feature characterising the contrast between these two pairs of stops. I argue that the contrasting feature between Czech alveolar /t, d/ and palatal /c, ɟ/ is the dependent feature [posterior]. Under the theory of segmental acquisition assumed in this thesis, if English requires the feature [posterior] for any phonemic contrast, then native English speakers should be able to perceive any non-native contrast of this feature, including Czech alveolar /t, d/ versus palatal /c, ɟ/. I will show that English contrasts three coronal fricative places of articulation: alveolar /s, z/ versus alveo-palatal /ʃ, ʒ/ versus (inter-)dental /θ, ð/ and requires the feature [posterior] as a dependent node of Coronal to contrast alveolar /s, z/ versus alveo-palatal /ʃ, ʒ/. Following Brown's theory, English speakers thus have the building block necessary for perception and eventual production of the novel Czech phonemic contrast of alveolar /t, d/ versus palatal /c, ɟ/. The results of perception and production experiments support this hypothesis.

The thesis is organised as follows. In Chapter One I outline the phonological assumptions held by Feature Geometry and Underspecification theories as well as the issue of the unmarked status of coronal segments. These theories provide the framework for the discussion of the acquisition of segmental phonology by first language learners discussed in Chapter Two. In Chapter Three we turn to second language learners' acquisition of new segmental contrasts. Chapter Four sets up the Feature structures of Czech and English, while in Chapter Five we turn to the experimental research conducted to test the model of L1 interference in L2 phoneme acquisition assumed in this thesis.

CHAPTER ONE

PHONOLOGICAL THEORY

0.0 BACKGROUND

In this chapter I will outline the phonological frameworks of Feature Geometry and Underspecification theory in order to lay the foundation for the generative model of second language segmental acquisition I have adopted in this thesis. In Section 1. we look at the theoretical framework of Feature Geometry theory as a formal means of representing structural relationships between features, the sub-units comprising individual speech segments. In Section 2. we look at the claims of Underspecification theory and the representation of redundant features. Finally, in Section 3 we discuss the marked status of coronal segments as well as arguments for the coronal place of articulation as cross-linguistically underspecified.

1.0 FEATURE GEOMETRY THEORY

Seminal work in phonological theory by Jakobson, Fant, & Halle (1952) provided an influential analysis of individual speech sounds in which individual sound segments (for example, the sound [b]) are claimed to be the sum of smaller sub-units, or features, rather than an indivisible entity. The motivation behind features stems from Prague School phonemic analysis, where segmental contrasts are dependent on contrast between individual features smaller than the segment (Trubetzkoy, 1939/1969). The claim that features, and not speech segments themselves, provide the basis for phonological contrast is supported by arguments that each feature represents an articulatory or acoustic component of speech production. Thus, two speech segments may have in common all features but one: this differing feature is sufficient to create phonemic contrast between segments.

Moreover, features enable us to group segments into natural classes. For example, languages may ban all voiced consonants from the coda position: the feature [voice] enables us to capture this generalisation whereas an analysis based on the

segmental level would not recognize this commonality and would result in a much weaker (less constrained) empirical hypothesis.

Early feature-based phonological theories, notably Chomsky & Halle's *Sound Pattern of English* (1968) (hereafter SPE), represented features in the form of an unordered, linear, feature matrix in which each segment was represented by either a positive or negative value for each (relevant) feature. For example, the representation of the phoneme /b/ as an unordered feature bundle would be as in (1):

$$(1) \quad \left[\begin{array}{l} +\text{consonantal} \\ -\text{syllabic} \\ -\text{sonorant} \\ +\text{anterior} \\ -\text{coronal} \\ -\text{continuant} \\ +\text{voice} \end{array} \right] = /b/$$

The disadvantage to this approach is that it does not capture natural relationships between features. SPE represented phonological processes as a series of rules operating when the environment corresponds to the representation. For example, the fact that English nasals assimilate to the place of articulation of the following consonant can be stated as the rule shown in (2):

$$(2) \quad \left[\begin{array}{l} +\text{nasal} \\ +\text{anterior} \\ -\text{back} \end{array} \right] \longrightarrow [-\text{coronal}] \quad / \quad ______ \quad [-\text{coronal}]$$

However, these types of rules describe rather than explain: the assimilation process characterised in (2) is equivalent to an articulatorily impossible process which assimilates

three arbitrary features.

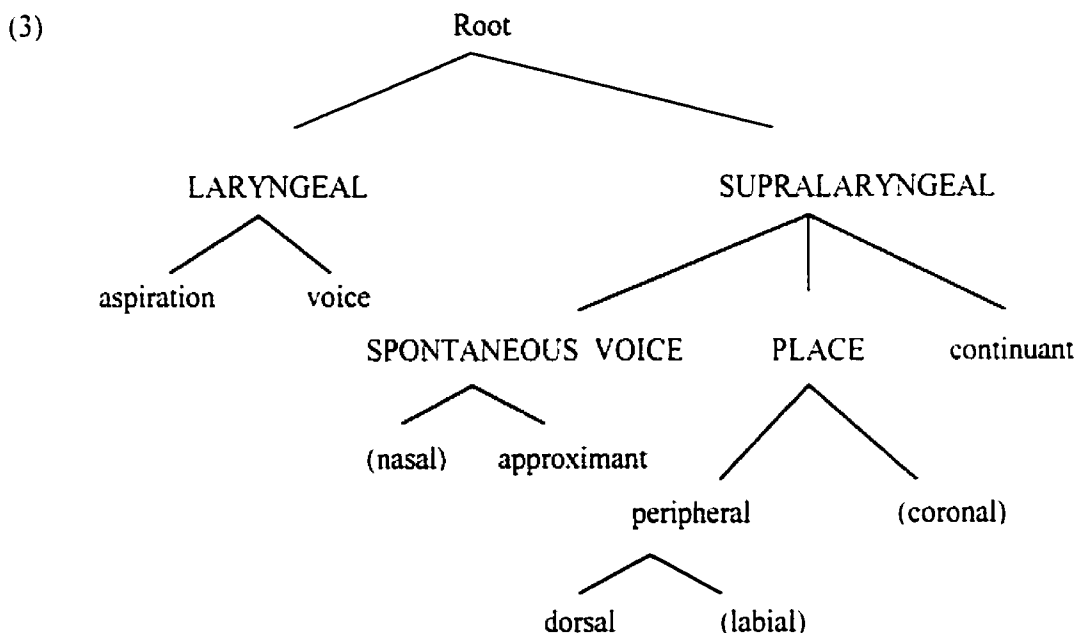
The understanding of segmental relationships and processes was greatly advanced with the proposal that unordered feature matrices be replaced with hierarchical structures in the framework of Feature Geometry (Clements, 1985; Sagey, 1986; McCarthy, 1988). Hierarchical feature geometries provide an advantage over linear models for several reasons.

First, as hierarchical models are based on the structure of the vocal tract they are better able to represent the fact that specific features tend to predictably co-occur with other features, thereby linking the physical phonetic implementation and the mental phonological representation.

Secondly, feature geometries can capture a range of diverse phonological operations with a cohesive set of defined operations such as Spreading, Delinking and Fusion (Avery & Rice, 1989). Thus, hierarchical models result in a more constrained theory as they provide a formal means of capturing natural classes of sounds via relationships of dependency and constituency between features.

To facilitate the discussion of hierarchical models and to provide a framework for the discussion of dependency and constituency relations, I present the model I assume in this thesis in (3)¹. This model is based on Brown (2000) which I have revised by placing peripheral as a dependent of the Place node, with secondary Content nodes below. This revision is based on child acquisition data discussed in Chapter Two.

¹ A precise hierarchical model and the features comprising it is still under some debate; however, the formal properties of various Feature geometries remain constant. The arguments presented in this thesis do not depend on the correctness of this particular model.



The hierarchical feature structure in (3) represents the sound segment as a cohesive unit gathered under and structured by the Root node. There are two types of nodes beneath, or dependent upon, the Root: (i). Organizing nodes, and (ii). Content nodes (Avery & Rice, 1989).

Organizing nodes (represented by means of all capital letters) serve to represent major organizational units based on the structure of the vocal tract as well as to define sets of features that pattern together with respect to phonological processes such as assimilation or spreading. The four major organizational units include the LARYNGEAL node, which describes states of the glottis (McCarthy, 1988); and the SUPRALARYNGEAL node, which subsumes both the PLACE node defining places of articulation such as labial, coronal, and dorsal, and the SPONTANEOUS VOICE (SV) node which distinguishes between sounds made with a lowered velum (nasal sounds) and those made with the velum raised (oral, or non-nasal, sounds).

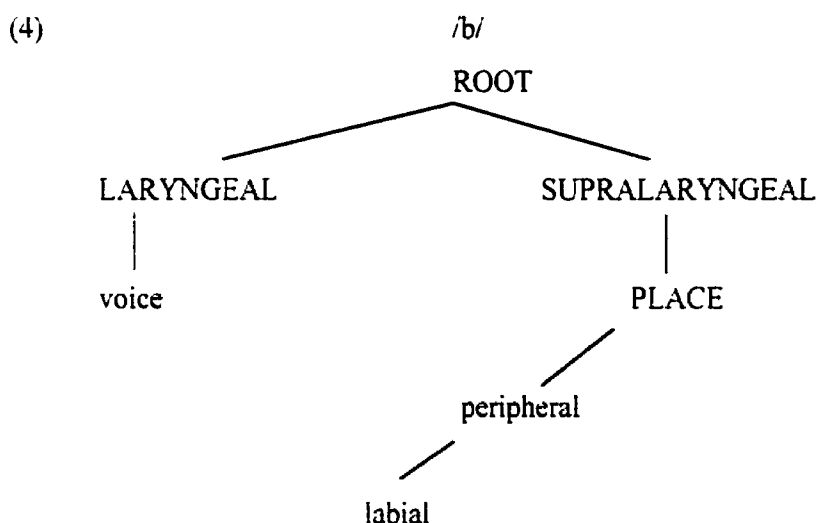
The second classification of nodes are the Content nodes, which define action of the articulators and are dependent on (or constituents of) the Organizing nodes. Content nodes occur in pairs of which one member of each pair is the default (or unmarked) node.

interpreted only in the absence of the other member. In (3), the features indicated between parentheses are the default features: they are not present in the underlying representation of a segment, rather, they are interpreted by the absence of the opposite feature (Brown, 1997). For example, the feature coronal under the Place node will be interpreted by default only if its opposite feature [peripheral] is not specified.

Content nodes are further subdivided into primary and secondary nodes: primary Content nodes are daughters of the main articulatory nodes and correspond to broadly defined movements of the articulators. Secondary content nodes are dependents of the primary nodes and provide fine-grained articulatory instructions. Let's look at a concrete example to illustrate: the Organizing node PLACE has as dependents the primary Content nodes coronal and peripheral to define broad articulatory gestures; these primary content nodes in turn have as dependents the secondary content nodes [anterior], [distributed] and [round] to provide detailed, fine-grained instructions to the articulators.

Crucially, the secondary nodes can be either redundant or distinctive in a language. When a feature is distinctive (that is to say, contrastive) in a particular language, it must be present in the structural representation. The case of English fricatives is illustrative: given that English has three fricatives produced at the coronal place of articulation, the alveolar /s/, the alveo-palatal /ʃ/ and the (inter)dental /θ/, the primary Content node coronal is not sufficient to distinguish between them and so finer articulatory detail must be provided by means of the secondary content nodes. Without this additional, finer grained structural information the contrast between the 3 sounds would be conflated.

Looking again at the sample segment /b/ that we saw linearly represented in figure (1), under a hierarchical representation the structure would be as in (4):



Comparing the linear structure in (1) with the hierarchical structure in (4) highlights another difference between these two representations: in structure (1), the segmental features are binary properties, where [+] indicates that a feature is present and [-] indicates absence of a feature, while in structure (4) the features are monovalent, or privative, and it is only the absence of non-relevant features that gives the appearance of binarity.

Phonologists disagree whether features are best represented as binary, unary, or some combination of the two, with the most vigorous debate engendered at the level of the secondary Content nodes such as [anterior] and [distributed]. This is partially due to the fact that precisely which features are required to characterise particular sounds in various languages is still under debate. Phonologists working within a Feature Geometry framework have argued for the use of privative features for some nodes, where only a single value (generally the marked value) is indicated. This means that generalisations can be made only of the class of sounds that possess the value since the group of segments that does not possess the feature do not together form a natural class.

Proponents of Articulator Theory (Steriade, 1987; Clements, 1988; Archangeli, 1988) claim that nodes corresponding to articulators such as Labial, Coronal and Dorsal are privative while the secondary content nodes such as [anterior] and [high] are binary: Avery & Rice (1989) argue that all features are monovalent. Van der Hulst (1989) also argues for monovalency as it is more restrictive: by representing features as monovalent,

the theory is more constrained since unary features incorporate relationships that would otherwise have to be listed as default rules applying later in the derivation. However, in claiming that features are monovalent we cannot then utilise the positive aspects of binary theory by replacing a single binary feature with two privative features. That is, we cannot posit two separate unary features [anterior] and [posterior] to replace [+anterior] and [-anterior]. However, as feature monovalency is a more restrictive theory, in this thesis I will assume that all features are monovalent and the presence of a feature in the representation of a segment indicates that the corresponding articulator is active: conversely, the absence of a feature indicates that the articulator is not active for that segment.

Another reason that the hierarchical structure of features shown in (3) is an advance over linear representations is that hierarchical representations better represent relationships of dependency and constituency between features, at all levels of the structure. Current phonological theory has taken the position that phonological processes are more elegantly explained via the representational component than by utilising rules (Yip, 1988; Piggot, 1988; Avery & Rice, 1989). Feature Geometric representations capture constituency relations in that Organising nodes represent articulatory movements and all features represented below the Organising node are constituents of that node.

As features capture natural classes, any process such as assimilation or spreading that affects a dominant node must necessarily characterise the subordinate node as well. For example, if a segment is specified for the feature [posterior], then it must necessarily be specified for the dominant node Coronal. The reverse, however, does not hold: if a segment is specified for the feature Coronal it is not also specified by default for the secondary feature [posterior]. The dependent feature [posterior] must be explicitly specified for a Coronal segment to be [posterior]. It is these processes of constituency and dependency that assist in specifying a language-specific hierarchy so that if a particular process cannot be shown to characterise all dependent segments of an organising node, then the proposed hierarchy is incorrect.

The claim for segments having internal, hierarchical structure, while allowing for more elegant elaborations of featural processes and relationships, has also raised questions as to precisely which features are present in the underlying phonological

representation, and which are absent to be added by default rule at the level of phonetic implementation (Steriade, 1987; Kiparsky, 1982). This brings us to the question of feature underspecification which I present in the following section.

2.0 UNDERSPECIFICATION THEORY

While theories of Feature Geometry are in agreement that individual speech segments have internal, hierarchically organised, structure, the precise representation of a particular segment is theory-specific. Most phonologists agree that redundant information need not be represented underlyingly, but can be added later in the derivation by rule. However, theories differ as to claims regarding the precise specification of redundant features which are predictable by the nature of the particular segment and are thus not necessary to indicate contrast. Redundancies may be either absolutes or determined by markedness (Ingram, 1995). Absolute redundancy can be seen in the case of vowels: if a vowel is specified for the feature [+high] it would be redundantly specified for the feature [-low] since no further information is gained by this specification.

Redundancies can also be specified by markedness: for example, since voiceless nasals and liquids are cross-linguistically very rare, they are argued to be marked in contrast to voiced nasals and liquids, thus, specifying liquids and nasals for the feature [voice] would be redundant. Crucially, as unspecified features are not assigned a value they are thus absent from the phonological representation. On the basis of this claim, underspecification theories argue that predictable or redundant features of a language do not need to be specified underlyingly, rather, they can be added later at the level of phonetic implementation². Only non-redundant features need be overtly specified since they are unpredictable and thus cannot be derived. Leaving predictable features out of the underlying representation has the effect of simplifying rules of assimilation. As Stemberger and Stoel-Gammon (1989: 182) note, underspecification is a method of building in a frequency bias: "The use of underspecification with a default feature-filling

² The theory of predictable feature underspecification has been challenged recently in the generative framework by constraint-based theories such as Optimality Theory (see Prince & Smolensky, 1993).

rule amounts to extracting the most frequent value of a feature for a given class of segments and building a bias into the language system to use that value of the feature unless it is specifically contradicted by other phonological information.”

Broadly, there are three main Underspecification theories distinguished on the basis of the level of representation required for redundant features: Contrastive Specification (Steriade, 1987; Clements, 1988) which requires all contrastive features, including redundant ones, to be overtly specified; Radical Underspecification (Archangeli, 1988; Paradis & Prunet, 1990, 1991) which eliminates all redundant features; and Minimally Contrastive Underspecification which takes the middle ground between the two earlier theories by specifying some redundant features while eliminating others on the basis of language-specific contrasts.

While differing in the level of representation required for redundant features, all three frameworks share three tenets: (i), the set of possible features is constrained or limited by Universal Grammar, (ii), no single language exploits all features (that is, each language uses a subset of the set of UG constrained features), and (iii), individual speech segments in a particular language are a subset of all the possible features of that language (a subset of the subset of features). In the next section we will look at each of these three theories in some detail.

2.1 *Contrastive Specification*³

The theory of Contrastive Specification determines underlying feature representations based on language-specific phonemic inventories. If two segments in a language contrast on the basis of a given feature, then both members of the contrasting pair must be specified for that feature with the corresponding positive or the negative value. If there is no segmental contrast, then specification is unnecessary. That is, if a particular feature is not necessary for contrast, then neither value is present in the underlying representation.

The assumption in Contrastive Specification is that overt specification is necessary to distinguish between two contrastive segments, but if there is no contrast then

³ Also known as Contrastive Underspecification

the feature value can be supplied by default rules based on markedness considerations. To illustrate, Paradis and Prunet (1991: 7) present a hypothetical language which contrasts three stops /p, b, g/, but lacks the voiceless velar counterpart /k/.⁴ In this case, /p/ is thus specified for [-voice] and /b/ is specified for [+voice], but /g/ is not specified for the feature [+voice] at all since there is no /k/ in contrast. Thus, language-specific segmental inventories provide the basis for representation of features in Contrastive Specification.

2.2 *Radical Underspecification*

Radical Underspecification takes the position of underspecifying redundant features a step further than does Contrastive Underspecification in claiming that most phonological features are redundant since they are predictable and can thus be specified according to universal markedness conditions. Radical Underspecification borrows from markedness theory in claiming that it is only the marked, or unpredictable, value that is present underlyingly since marked features are language-specific and cannot thus be predicted. Unmarked features provide the default values and thus need not be specified as they can be filled in by default rules in the phonetic component. Unlike Contrastive Underspecification which requires both values of a contrastive feature to be represented, Radical Underspecification requires only a single value to be specified. Thus Radical Underspecification makes different claims for non-contrastive phonemes than does Contrastive Underspecification. Returning to the hypothetical language presented by Paradis and Prunet to illustrate: if a language has the three stop segments /p, b, g/, but not /k/, and we assume that the unmarked value for voicing is [voiceless], then under Radical Underspecification the voiceless segment /p/ will be unspecified for the default feature [voice] while both /b/ and /g/ are underlyingly specified as [voice] as this value is unpredictable.

⁴ This example is problematic: as voiceless segments are cross-linguistically unmarked, it is unclear how markedness considerations would yield voiced /g/ and not voiceless /k/.

2.3 *Minimally Contrastive Underspecification*⁵

In response to problems encountered by Radical Underspecification and Contrastive Specification, Avery & Rice (1989) propose a modified version of Underspecification theory called Minimally Contrastive Underspecification (MCU), which arrives at a middle ground between the two theories by borrowing key elements from both. As in Contrastive Underspecification, the key to the representation of features in Minimally Contrastive Underspecification is based on language-specific phonemic inventories. Both Contrastive Underspecification and MCU claim that the structural representations of contrasting segments are specified with a particular feature when that feature is necessary to maintain contrast between segments, but if a feature is redundant, it need not be specified. Because phonemic contrasts vary from language to language, featural representations are also language-specific. As with Radical Underspecification, it is the marked value of a feature that is specified when needed to maintain a contrast between segments.

However, Minimally Contrastive Underspecification differs from Radical Underspecification and Contrastive Specification in that contrasts between segments can *trigger* the specification of a node where that node would be underspecified or derived in Radical Underspecification, and present underlyingly in Contrastive Specification. This notion of node triggering is formulated in the Node Activation Condition given in (5):

(5) Node Activation Condition (NAC)

If a secondary content node is the sole distinguishing feature between two segments, then the primary feature is activated for the segments distinguished. Active nodes must be present in underlying representation.

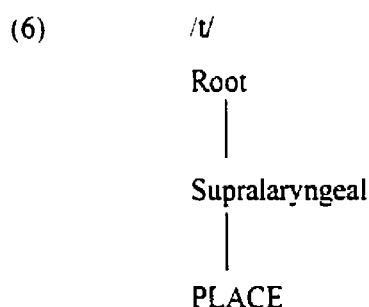
(Avery & Rice, 1989:183)

The Node Activation Condition can be illustrated by looking at the representation of coronal segments. The NAC holds that if a particular language has only a single coronal place of articulation, say the alveolar /t/, then the Place feature [coronal] as the unmarked

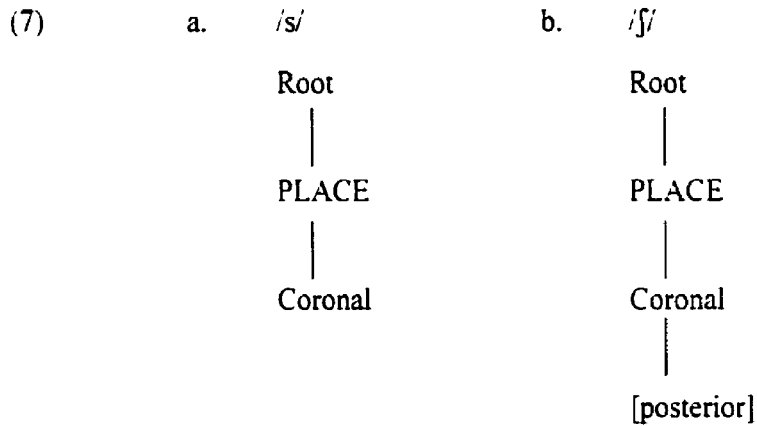
⁵ This theory is also known as Minimally Contrastive Specification

value is underspecified, or not represented. The single coronal segment is represented by a bare Place node. However, if a language has more than one sound articulated at the coronal place of articulation, then one of the segments must be represented by means of a secondary node such as [anterior] to distinguish between the two coronals. The representation of the secondary node [anterior] triggers the specification of the primary node Coronal for those segments which contrast on the basis of this secondary node. Let's look at English to illustrate.

English has a single coronal place of articulation for stops, the alveolar /t/. Because there is no other coronal segment in contrast, the Coronal node remains underspecified and /t/ is represented by a bare Place node. The representation of /t/ is shown in (6):



However, English has two fricatives in contrast under the coronal place of articulation: an anterior coronal, the alveolar /s/; and a posterior coronal, the alveo-palatal /ʃ/. Representing two coronal fricatives solely by means of the feature Coronal would not provide enough articulatory information to distinguish between them, so further elaboration of the coronal node is required by means of a secondary node. Under the Node Activation Condition, if a single secondary content node is the distinguishing feature between two segments, then the primary feature is activated for these segments. In the case of English fricatives, the specification of the secondary content node [posterior] for the segment /ʃ/ triggers the specification of the dominant coronal node for both segments. The representations for /s/ and /ʃ/ are shown in (7):



Under the framework of Minimally Contrastive Underspecification, then, featural specification is dependent on both language-specific phonemic contrasts as well as markedness considerations. This framework underlies Brown's theory of phoneme acquisition, which I have adopted in this thesis.

In sum, Underspecification theories make claims as to which features need be specified underlyingly, and which can be derived at the level of phonetic implementation. In the next section, we look at arguments for the coronal place of articulation as the universally or cross-linguistically underspecified place of articulation.

3.0 CORONAL UNDERSPECIFICATION

Coronal segments are those sounds that are produced with the front part of the tongue, which encompasses the tongue tip and tongue blade, and include five primary places of articulation: dental, alveolar, palato-alveolar, retroflex, and palatal⁶ (Maddieson, 1984). Phonetically, the five places of coronal articulation are distinguished from five other main places of articulation, that is, bilabial, labiodental, velar, uvular, and pharyngeal, which means that coronals make up half of the primary places of articulation. (Keating, 1991). Ladefoged (1982) points out that the tongue tip and blade are the most mobile parts of the tongue, and the tongue blade is conducive to a greater variety of

⁶ Ladefoged & Maddieson (1988) include two less common coronal places, linguolabial and interdental.

articulations than are other articulators. I will provide a more detailed description of alveolar and palatal sounds, the two coronal places of articulation relevant to this thesis, in Chapter Three.

3.1 *The Special Status of Coronals*

Cross-linguistically, the coronal place of articulation is claimed to be the unmarked or default place of articulation, based on several characteristics differentiating them from segments produced with the lips (labial sounds) or those produced with the tongue body (dorsal sounds). Stemming from seminal work by Kean (1975), many phonologists (see for example the articles in Paradis & Prunet, 1991) have argued for coronals as the most neutral or unmarked place of articulation based on three unique properties. First, coronal sounds are cross-linguistically very frequent, in that all languages with the possible exception of Hawaiian include at least one coronal stop in their inventory (Maddieson, 1987: 31). In a survey of 317 languages, Maddieson found that 316 have the coronal nasal /n/ in their inventory; liquid sounds are coronal in the majority of languages; and if a language has only one fricative segment, then it will be the coronal /s/ 84% of the time.

Secondly, coronals appear to play a unique role in child phonology. Both coronal stops and fricatives (along with labials) are among the first consonants to be acquired by children (Stoel-Gammon, 1985; Vihman, Ferguson, & Elbert, 1986). Moreover, harmony processes in child acquisition also point to the special status of coronals. Consonant harmony involving two non-adjacent consonants is a common process in child language as shown in (8):

- (8) a. [gʌk] ‘duck’ (velar harmony)
 b. [bʊp] ‘boot’ (labial harmony)

Stemberger & Stoel-Gammon (1989) point out that harmony processes in assimilatory contexts tend to replace underspecified elements with specified ones so that children will tend to assimilate alveolars to velars and labials, but not the other way around.

Thirdly, coronals are characterised by phonological processes that do not affect other places of articulation, such as assimilation, neutralization, and transparency. Kiparsky (1985) points out that coronal sounds are more likely to undergo assimilation of place features than are labial or velar consonants. In Catalan, for example, Kiparsky argues that the fact that the coronal nasal /n/ in Catalan assimilates to all consonants while the labial /m/ assimilates only to labiodentals while retroflex and velar nasals /ɲ, ŋ/ do not assimilate at all can be accounted for by underspecifying the Place features for coronal. If the Coronal place is underspecified, then the assimilation of coronal nasals is a result of the spreading of Place features from the following consonant to the underspecified coronal.

Another process that characterises coronals is transparency, where a segment allows a feature to spread across it as a result of being unspecified, or transparent. Specified segments are opaque to spreading. As Paradis & Prunet (1991) point out, if coronals are underspecified for Place in the underlying representation, then transparency effect should single out coronals. They claim that the transparency effects found in some West African languages whereby vowel spreading is blocked across non-coronal segments but not across coronals can be accounted for if it is the Place node of the vowel that is spread to the underspecified Place node of the Coronal. Spreading would thus be blocked by non-coronal segments as they are already specified for Place.

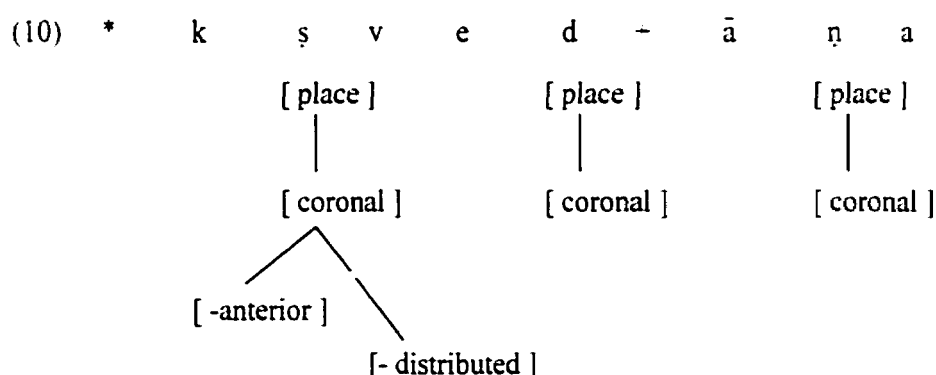
3.2 *The Underspecification of Coronals*

Because coronal sounds appear to have a special status among segments, the Coronal place of articulation is argued to be, cross-linguistically, the default or underspecified articulator. Under Articulator Theory, Coronal is classified as an Articulator node, along with Labial and Dorsal. Nodes corresponding to articulators are privative so that the feature [-coronal] cannot exist. Most phonologists working within the framework of feature geometry argue for [coronal] as a privative feature (Sagey, 1986; McCarthy, 1988; Yip, 1989; Clements & Hume, 1995). By claiming [coronal] as a unary feature, then coronal sounds are marked for the presence of this feature, while non-coronal sounds are not. This has the effect that non-coronals cannot be grouped together as a natural class as they could under a binary framework unless dominated by

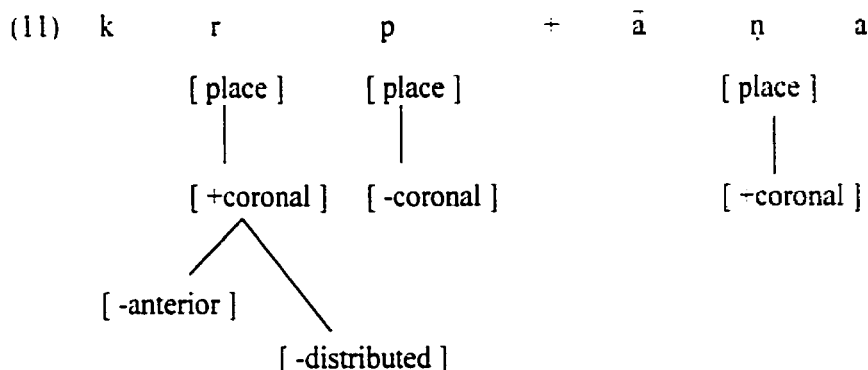
the Peripheral node grouping labials and dorsals together as a natural class. As Roca (1994) points out, eliminating the feature [-coronal] is desirable if we want to account for cases such as the well known process of *n*-retroflexion in Sanskrit (Whitney, 1885; Odden, 1978), where the alveolar [n] surfaces as retroflex [ɳ] if it follows a retroflex [ʂ] or [ɭ] without a coronal segment intervening. This process can take place at some distance, with as many as 4 non-coronal vowel or consonantal segments intervening between the trigger and the target. That is, rules spreading [+coronal] are blocked by [coronal] sounds but not by labials and velars: assimilation cannot occur across an opaque coronal segment. In Sanskrit, the alveolar [n] is retroflexed to [ɳ] if it follows a retroflex [ʂ] or [ɭ]: This is illustrated in (9):

- (9) a. kṣubh + āṇa 'quake'
 b. kṣved + ana 'lament'

The representation of (9b) is shown in (10):



The blocking effects can be seen as a violation of the line-crossing constraint, where the intervening coronal /d/ blocks the spreading of the feature [coronal]. Roca argues that the transparency of non-coronal segments to the spreading of the feature [coronal] cannot be explained if labials and velars are specified as [-coronal], since the feature [-coronal] would be on the same tier as [+coronal]. This would have the effect of incorrectly blocking the spreading of [+coronal], as illustrated in (11):

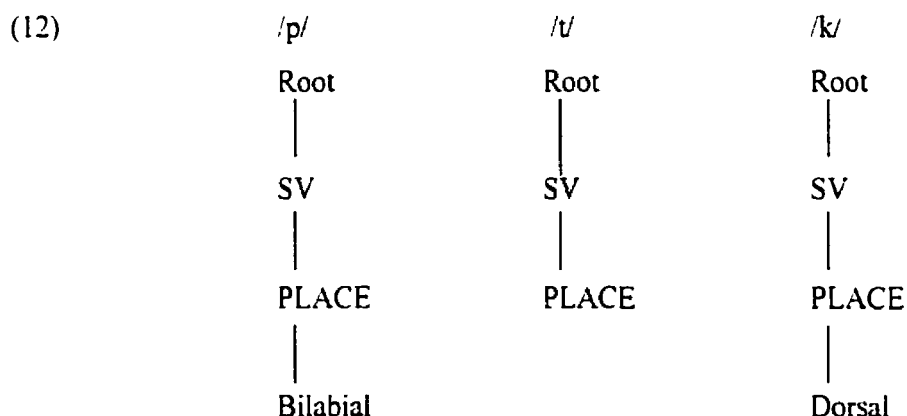


A number of phonologists have argued for underspecification of the coronal node as an alternative to using [-coronal] so that coronals are not specified for a Place node at all, while non-coronals are specified for Place: specifically, labial sounds such as bilabials and labiodentals represent Labial as a dependent of the Place node while velar, uvular and pharyngeal sounds are specified for Dorsal as a dependent of the Place node. If there is contrast within the class of coronals, then specification of the coronal node is subsequently triggered by the Node Activation Condition.

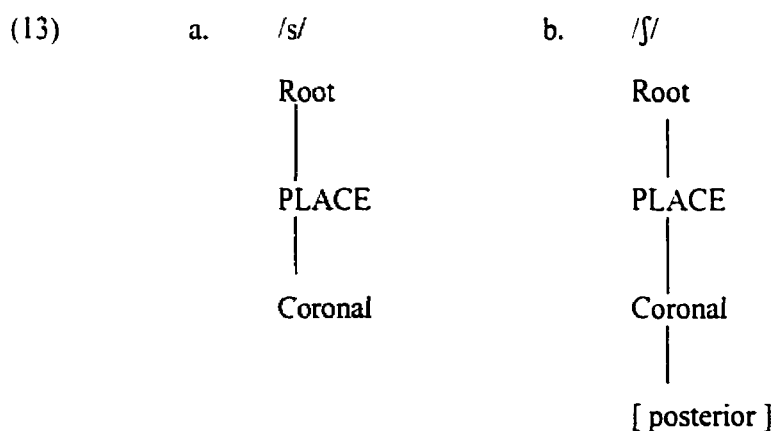
Support for Underspecification of Coronal segments for Place features is provided by Yip (1991) and Paradis & Prunet (1989), among others. Yip argues that the nonoccurrence in English of word final and medial clusters (such as pk, kp, fk, mk) which contain more than one non-coronal segment can be attributed to Underspecification of the Place node for Coronals. Yip formulates this insight as a Condition on Clusters, which states that adjacent consonants are limited to at most one Place specification. (Yip, 1991:62). This filter blocks non adjacent coronals with more than one coronal because labial and velar sounds are specified for Place. However, two-coronal sequences such as [st], [lt] and [nt] are allowed word finally in English because coronals lack specification for Place.

Underspecification of the Coronal articulator node means that if a language has only one segment produced at the coronal place of articulation, as the unmarked place of articulation it need not be specified. The interpretation of the Coronal Place node is supplied by default rules at the level of phonetic implementation. Again, this can be

exemplified by English, which has a single stop consonant articulated at the Coronal place of articulation: the alveolar stop /t/. Thus, the representation for coronal /t/ has a bare unspecified Place node since further elaboration of the coronal node is not necessary for contrast within the class of coronals⁷. This is presented in (12):



Coronal segments are not represented for place of articulation in the underlying representation unless that particular language has phonemic contrast *within* class the of coronal segments. As we saw in Section 2.3 above, this is the situation we find in the case of the English coronal fricatives where alveolar /s/ and alveo-palatal /ʃ/ contrast. The representations for these two coronals are repeated here as (13):



⁷ The representation abstracts away from other nodes such as Laryngeal and Air flow.

Under the Node Activation Condition, the Coronal Place specification is also triggered because the secondary Content node [posterior] is required to distinguish between the segments /s/ and /ʃ/.

In sum: in this chapter we looked at the mental representations of phonological segments and processes in the frameworks of Feature Geometry and various Underspecification theories. In this thesis I assume the framework of Minimally Contrastive Underspecification, as this approach to Underspecification is dependent on both language-specific phonological contrasts as well as markedness considerations. This chapter also discussed arguments for the status of Coronal segments as the cross-linguistically unmarked, or default, articulator. I assume that the Coronal node is unspecified, with secondary or dependent nodes present only when more than one sound is contrasted at the coronal place of articulation. Discussion of the arguments and terminology of these two theories provided the necessary background information for the discussion of first language segmental acquisition in Chapter Two, as well as second language acquisition in Chapter Three.

CHAPTER TWO

FIRST LANGUAGE ACQUISITION OF SEGMENTAL PHONOLOGY

0.0 BACKGROUND

In Chapter One, we looked at the phonological assumptions held by Feature Geometry and Underspecification theories. These phonological theories provide the framework for the discussion in Chapter Two on the acquisition of segmental phonology by first language learners. In turn, this chapter on first language phoneme acquisition provides a background to theoretical assumptions regarding second language learners' acquisition of new segmental contrasts to be discussed in Chapter Three. In order to understand what adult language learners must acquire in learning the sound system of a second language, it is essential to understand the developmental process that children undergo when acquiring the phonological contrasts of a first language.

Chapter Two is structured as follows: Section 1 provides a general overview of the complex task of first language acquisition as a whole. Section 2 moves from the general to the specific in looking at first language acquisition of segmental structure. In particular, we will look at Rice & Avery's (1995) arguments for the acquisition of speech segments as a process of structural elaboration, as well as Cynthia Brown's work (1993, 1997, 2000; also Brown and Matthews: 1993, 1997) integrating infant speech perception and the acquisition of phonological structures. This work in L1 perception and phonological acquisition forms the basis of Brown's model of L1 interference in second language phonological acquisition presented in Chapter Three.

1.0 FIRST LANGUAGE ACQUISITION

Before we turn to the acquisition of phonemes by children learning language, let us first take a look at child language acquisition in a more general sense. Language acquisition is a striking example of a conundrum that has challenged philosophers from Plato to the present: How are humans able to acquire such rich and complex systems of knowledge, given the relatively limited and fragmentary input they are exposed to? The

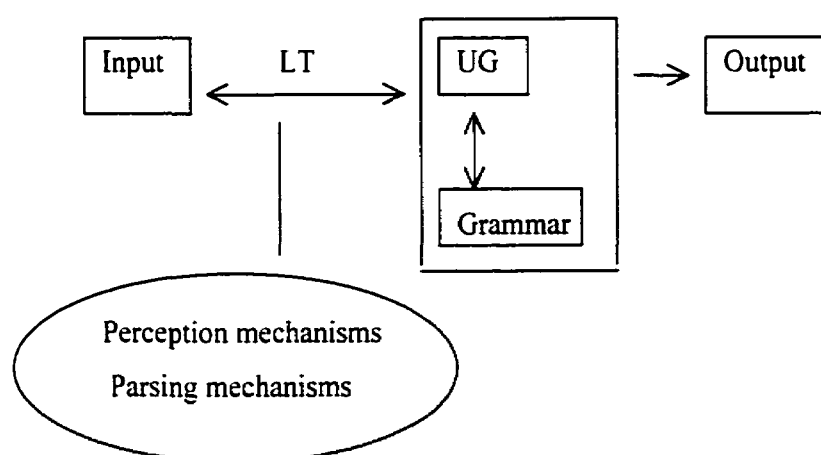
puzzle of how humans acquire the complex linguistic structures of language despite the comparatively impoverished input received has been called the “logical problem of language acquisition” (Lightfoot, 1981).

This gap between the rapid acquisition of complex grammatical structure and deficient input has led to the positing of an innate mental structure known as Universal Grammar, or UG (Chomsky, 1975), which at once constrains cross-linguistic variation while informing the process of language acquisition. Because UG is argued to be a module specific to language (as opposed to, say, memory) that works on linguistic input to produce a mental grammar, it falls within a rationalist theory of learning in which acquisition is a deductive process of moving from general, known principles to particular, language-specific grammars.

UG offers a solution to the logical problem of language acquisition by postulating an innate set of cognitive principles and parameters, where principles are broad universal structures underlying all languages, and parameters are language-specific characteristics triggered (in the manner of an on-off switch) by the input (Chomsky, 1988). This innate system of linguistic universals greatly simplifies the task of language acquisition in that the child comes equipped with this set of universal structures that do not therefore need to be overtly learned. By limiting the range of possible grammars, UG can account for the speed with which children acquire the complex grammatical structures of language. First language acquisition is thus a process of the innate, UG-provided, principles and parameters interacting with the language-specific input the child is exposed to.

The broad framework of acquisition that I will be assuming in this thesis is a modular model based on the notions of UG. The framework is shown in Figure 2.1, with an explanation of the individual terms following:

Figure 2.1: Modular Framework of Language Acquisition



In this modular framework, it is assumed that Universal Grammar (UG) is a language-specific module that works on linguistic input to produce a mental grammar and provides the child with much information regarding the underlying linguistic structure. In the domain of phonology, for example, the child may not need to overtly learn that all sound segments require Organizing nodes such as LARYNGEAL, AIRFLOW, SONORANT VOICE and PLACE, as this information is provided innately by UG; however, the child would have to learn which dependent features beneath the Organizing nodes are triggered by the linguistic input as these features vary from language to language.

INPUT is the ambient language provided by both caregivers and non-caregivers from which the child is supplied with phonological cues as to phonological constituents of that language. Through the inter-relationship between UG and the INPUT, the child arrives at a language-specific GRAMMAR, which is a mental representation of the target language.

Finally, we require some sort of learning theory (LT) mediating between the INPUT and UG to explain the sequence of grammars that the learner goes through in the process of acquiring the adult grammar. The most commonly held view in a modular theory is the system of principles and parameters (Chomsky, 1981, 1988). In this framework, input is viewed as a trigger where underspecified or default principles are

provided innately by UG, and the marked setting of the parameters must be set on the basis of experience with the target language input. Let us take a look at the acquisition of phonological stress to illustrate.

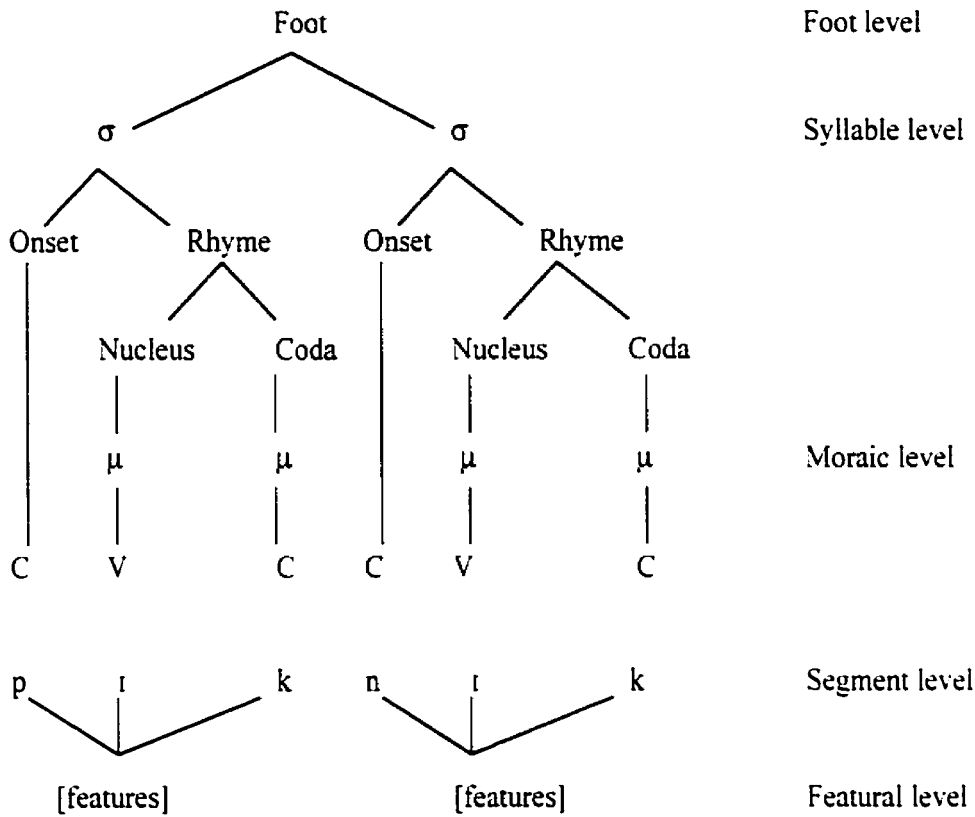
Metrical stress theory argues that stress is a manifestation of vowel or syllable prominences that are organised into prosodic units including the foot, word, and phrase. These prosodic units can be assembled in a constrained number of ways on a language-specific basis following a set of parameters. For example, Dresher & Kaye (1990) proposes a series of binary parameters for metrical stress, each parameter associated with and triggered by a particular phonological cue. Sample parameters are presented in (1):

- (1) a. Feet are [Binary/Unbounded].
- b. Feet are built from the [Left/Right].
- c. Feet are quantity-sensitive (QS) [Yes/No].

The general principles provided by UG constrain the metrical structures themselves. The learner's task is to determine how the parameters are set in the target language based on the ambient input.

Finally, the framework adopted in this thesis assumes both a Perception and a Parsing mechanism. As will be discussed in some detail in Section 2.2.1, research has shown that very young infants have the ability to perceive relevant aspects of speech such as pitch and voicing distinctions that will be useful to them when they begin to produce speech. Following work in syntax (see Fodor, 1999) on the claims for an innate parser to assign syntactic structure to a word string, work in phonology has similarly assumed a parsing mechanism to assign hierarchical structure to a string of sounds (see Dresher, 1999). A common model of phonological structure assumes the hierarchical levels shown figure 2.2 below:

Figure 2.2 *Model of Phonological Structure*



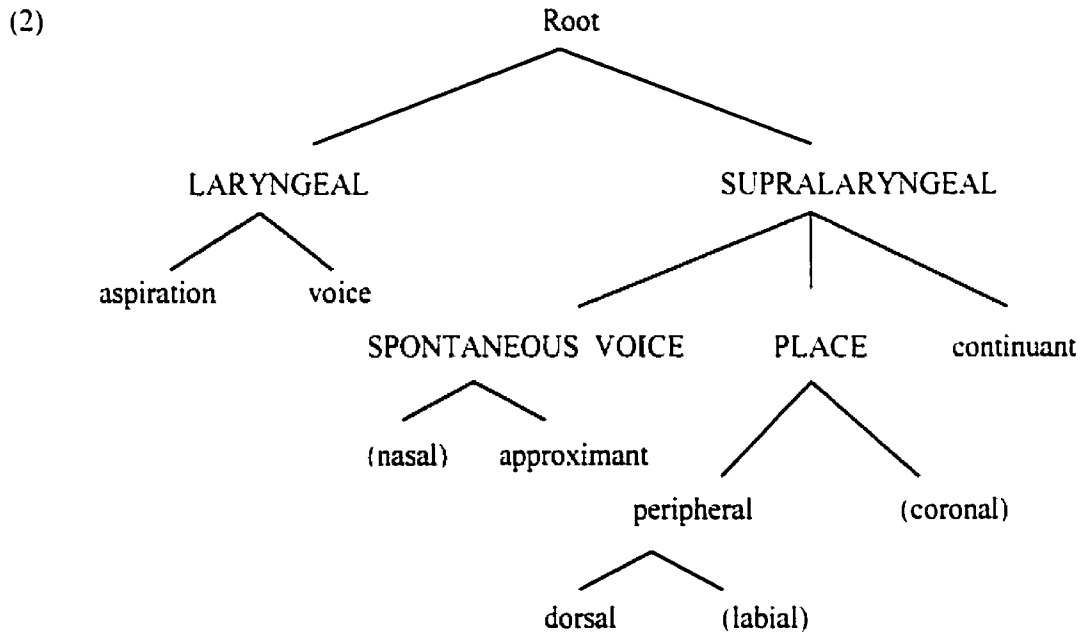
Now that we have looked at the puzzle of language acquisition in a broad sense, it is time to narrow our focus and turn to the area of how individual speech segments of a language are acquired by children. In the next section we look at the acquisition of hierarchical Feature Geometries as well as acquisitional models that have been proposed by Rice & Avery (1995) and Brown (1997, 1999).

2.0 ACQUISITION OF THE FEATURE GEOMETRY

The phonological component of language can be subdivided into two broad sections: (i), *segmental phonology*, which is concerned with the patterns and processes of phonological segments such as consonants and vowels; and (ii), *suprasegmental, or prosodic, phonology* which involves areas at a level above the individual segment such as

the syllable, stress assignment, and intonation. As we saw in Chapter 1, Feature Geometry theory provides a framework for explanation in segmental phonology by claiming that individual speech segments are composed of hierarchically structured sub-units called features. Within Feature Geometry theory there are two possible approaches to the acquisition of segmental representations: the theory of Full Specification versus the theory of Minimal Specification¹.

Full Specification claims that UG provides a fully elaborated feature structure representing all possible phonological contrasts. Redundant features are pruned away from the structural representation on a language-specific basis as the child becomes aware that not all possible phonemic contrasts are present in the ambient language. In this view, all children would initially have the fully elaborated feature tree shown in (2):

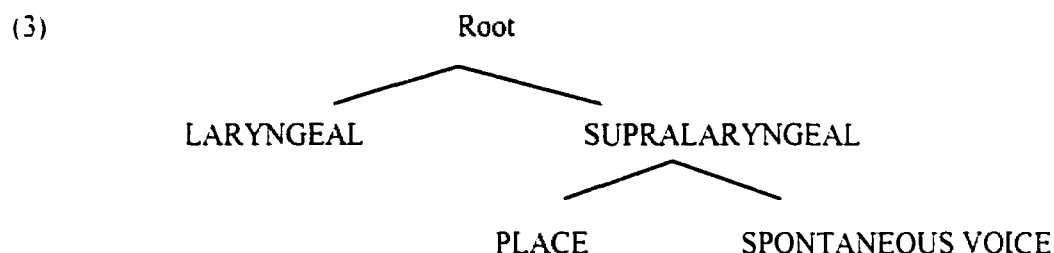


The disadvantage of the Full Specification approach, however, is its reliance on

¹ The theories of Full Specification and Minimal Specification are also known as the Pruning Hypothesis and the Building Hypothesis, respectively.

negative evidence: features are pared away only when a featural contrast is either redundant or not relevant to the target language.

In contrast, the theory of Minimal Specification takes the opposing stance in arguing for initial minimal structure, with features added to the structure in response to language-specific phonemic contrasts detected in the input. Once a child can productively contrast a particular phoneme, the features represented are only those which keep a segment minimally distinct from other segments in the inventory. Featural representations, then, are language-specific based on contrasts in the inventory. Under this approach, the acquisition of segmental contrast is a process of structural elaboration based on relationships of node constituency and feature dependency guided by the structural hierarchy. Thus, the goal of the language learner is to expand the language-specific structure until all the features that are necessary to distinguish all phonemes in the ambient language are present. In this view, children initially have minimal structure (representing only the Organizing nodes); further features are added in a step-by step fashion as needed in response to contrasts detected in the ambient input. The initial minimal structure is shown in (3) :



As will be discussed in upcoming sections, recent research in segmental acquisition (e.g. Brown, 1993; Brown & Matthews, 1997; Rice & Avery, 1995) argues for minimal specification by claiming that L1 acquisition of the phonemic inventory is a process of structure building in response to language-specific contrasts detected in the input, until all (and only) those features differentiating speech segments in a particular language are present. Based on experimental results, these researchers argue that the goal of the language learner is to expand the initial, UG-provided, minimal structure until all

the features that are necessary to distinguish all phonemes in the native language are present.

As this thesis assumes structural elaboration as opposed to pruning, the following sections are devoted to the arguments put forth by two main proponents of the structure building approach to segmental acquisition.

2.1 *Rice & Avery (1995)*

Recent models of generative phonology are also inherently models of phonological acquisition in that children's constructions of phonological representations are constrained and guided by principles of UG. Central to this assumption is the *Continuity Hypothesis*, which claims that although children's grammars are constantly evolving, each developmental stage conforms to universal linguistic principles (see for e.g. Lust, 1994, for the Strong Continuity Hypothesis; Clahsen, Eisenbeiss, & Vainikka, 1994; Paradis & Genesee, 1997, for the Weak Continuity Hypothesis). While a child's developing grammar may differ at a particular stage from the target adult grammar, at no stage in the acquisition process will the grammar violate UG principles. That is, the grammar may diverge (in some cases, substantially) from the target adult grammar, but being UG constrained it could be a possible grammar for some language.

Rice & Avery's (1995) model of phonological acquisition as segmental elaboration assumes continuity in arguing that each developmental stage of children's grammar is constrained by UG. Their model capitalises on the theoretical power of hierarchical models in arguing that acquisition of a phonemic contrast requires acquiring the language-specific featural structure characterising the two phonemes. As we saw in Chapter 1, the possible inventory of segmental features is constrained by UG, and no language utilises all the possible features or feature hierarchies, using instead only a subset of all the possible features. Moreover, each segment uses only those particular features required for contrast, not all features in the language-specific hierarchy. Because languages differ with respect to which segments are contrastive, speakers of different languages will acquire different structures, albeit constrained by Universal Grammar.

A crucial assumption held by Rice & Avery, and one that is adopted in this thesis, is that the order of phonemic acquisition falls out from the hierarchy of features so that

the acquisition of segmental structure occurs by a process of expansion of the language-specific feature hierarchy². Thus, the acquisition of new contrasts is a process of structure building, or elaboration, as opposed to the pruning of redundant features from an initial, fully specified representation. Moreover, elaboration of the feature tree is not a random process, but is instead guided by two principles consistent with the hierarchical structure.

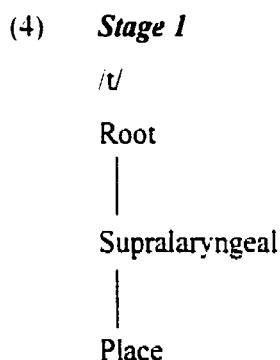
The first principle is the *Principle of Minimality* which holds that initially, the child has minimal featural structure. UG provides the emerging grammar with minimal structure which is then elaborated on a language-specific basis in response to detection of phonemic contrasts in the input. Crucially, this principle forces the assumption that initial phonological representations are impoverished, and that structure is permitted only in response to contrasts in the inventory (Rice, 1996). It should be emphasized that detection of a phonemic contrast is essential: the mere presence of a contrast in the input is not sufficient impetus for structure building.

The second guiding principle for segmental elaboration is the *Principle of Monotonicity*, which holds that feature inventories are built up in a monotonic, or step-by-step, fashion. New segmental structure is built in a node-by-node fashion based on the hierarchical relationships of constituency and dependency. Any intermediate structure posited by the child between the initial minimal structure and the fully elaborated target grammar will respect the hierarchical structure. Children therefore do not produce 'wild' grammars that do not conform to principles of UG. Thus, once a dependent feature has been acquired, it implies that the feature's superordinate node has already been acquired. Because a dependent feature cannot be acquired before a superordinate node, the implication is that children will be able to contrast those segments that have less structure before segments that are more structurally complex.

By way of illustrating the Principles of Monotonicity and Minimality and their implications for language acquisition, I present Rice & Avery's (1995) three-stage developmental path outlining the acquisition of Place distinctions characterising Labial.

² See Jakobson (1941/1968) for an earlier view of language acquisition as a process of increasing structure complexity.

Coronal and Velar sounds. As stated by the Principle of Monotonicity, acquisition of segmental structure proceeds by the elaboration of a single node at a time, or monotonically, following the pathways dictated by the hierarchical structure. Thus, if a language has only a single place of articulation, then only the dependent Place node is required to represent the unmarked default articulator, the coronal place of articulation. This is shown in (4):

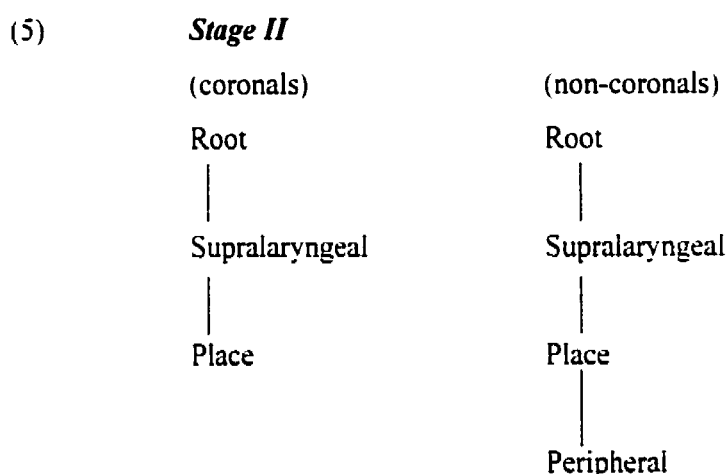


At this very early stage, the child has a single (unmarked) mental representation for /t/ without fine-grained place distinctions and can contrast a single consonant and a single vowel. The fact that the initial consonant produced by children is often [p], and not [t] as would be expected, is argued by Rice and Avery to result from a lack of motor control (for example, little control over tongue musculature) on the child's part: it is *not* the result of initial phonological specification as a labial sound³. At this stage of acquisition, the phonemes /p/ and /t/ would not be in contrast, although we may find both [p] and [t] produced as phonetic variants of /t/. Variability is linked to minimal structure: because little structure is specified at this point, the result is a broad phonetic range. Regardless of the exact phonetic implementation, however, this initial sound is unmarked. The initial lack of specification results in greater variability so that the unmarked sound may be realised as several phonetic variants of the single phoneme. The crucial prediction is that the first contrast in place of articulation will be between a coronal and a peripheral

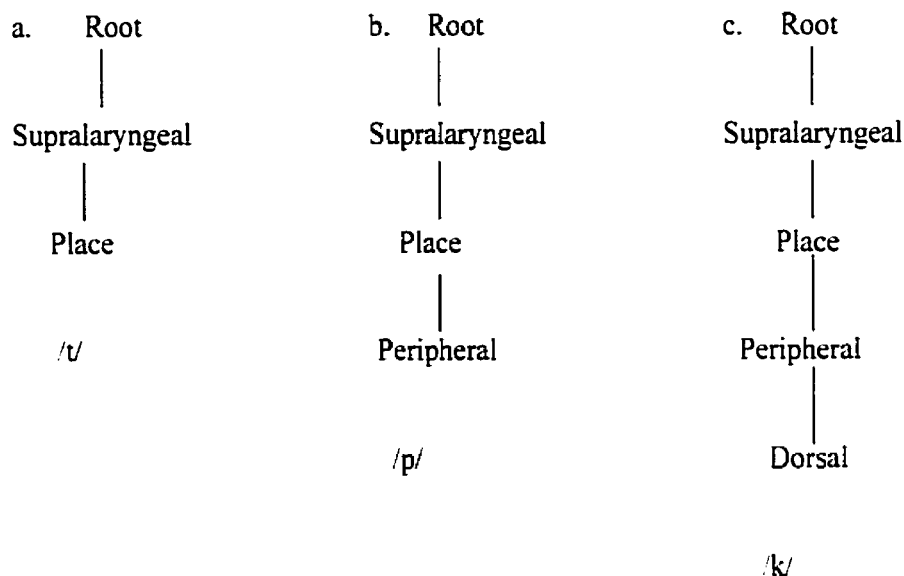
³ Jakobson (1941/1968) claims that the first consonantal sound is often /p/ in contrast with the vowel /a/ as these two sounds are "maximally distinct".

sound such as /t/ versus /p/.

If the language contrasts more than one place of articulation, then further elaboration of the Place node is required by the addition of the peripheral node. At this second stage, the child can contrast coronal and non-coronal sounds for a two way Place distinction, but does not yet phonologically distinguish *within* the class of non-coronals, that is, between labials and dorsals. That is, the child would not productively contrast /p/ and /k/. The elaborated structure for a two-way Place of articulation contrast is presented in figure (5):



At Stage III the peripheral node is expanded with the addition of the Dorsal node, creating a three-way place distinction with contrast between coronal and non-coronal (peripheral) sounds, as well as within the class of peripheral sounds (labial versus dorsal). Three places of articulation for stops are now contrastive: the child is able to contrast coronal /d/ versus labial /b/ as well as coronal /d/ versus velar /g/. Under the principle of Monotonicity, the Peripheral node must be acquired before its dependent Dorsal node. That is, no child would have the structure represented in (6c) below without first acquiring the representation in (6b). The hierarchical representation for Stage III is illustrated in (6):

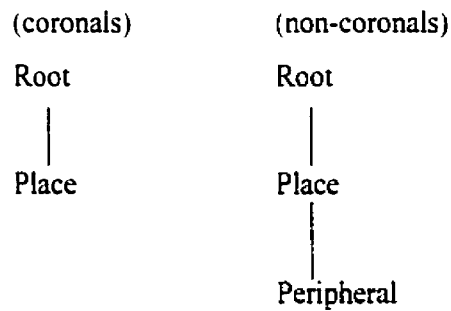
(6) *Stage III*

As we have seen, the theory of structural elaboration based on initial minimal structure can account for universal order of acquisition: it can account for the slight inter-language variability commonly found in child speech as well. Cross linguistically, children acquire particular segments before others, yet within this order there is variation. Rice (1996) points out that when children do not have a productive contrast between two segments, the amount of variability in the production of the sounds is greater than when the featural contrast has been acquired. In this framework, structural variation is a therefore a consequence of pathways within the hierarchical structure. Let's look at a concrete example to illustrate.

Within the Feature Geometric hierarchy, the organizing Supralaryngeal node subsumes both the Place and the Sonorant Voice nodes: this has implications for the possible order of acquisition of segmental contrasts represented by these nodes. Once a child has acquired the Supralaryngeal node, he or she has freedom to elaborate between the two dependent branches of the structure, expanding either the Place node or the Sonorant Voice (SV) node first. Thus, one child may elaborate the Place node before the SV node, producing contrast between the coronal Place of articulation and peripheral (or non-coronal) places of articulation. Since the SV node has not yet been elaborated at this

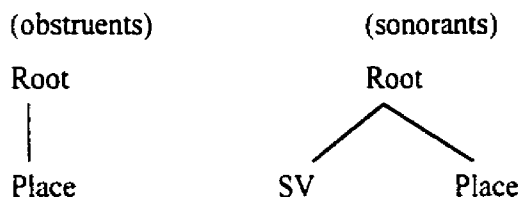
stage, the contrastive distinction between sonorants and obstruents is conflated. Figure (7) shows the representations for this initial stage with the elaboration of Place by the feature Peripheral:

(7) *Stage 1: Addition of the Peripheral node*



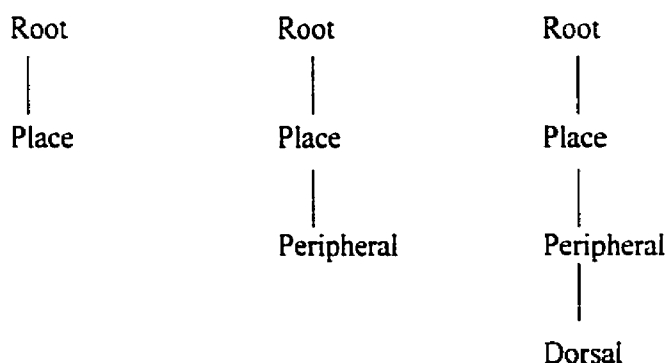
Both coronal and non-coronal consonants may be realised either as oral or nasal sounds -- the distinction is conflated -- because the SV node distinguishing between these two groups of sounds has not yet been elaborated. Moreover, because the Peripheral node itself does not have additional structure, the non-coronal segments are Labial by default, since the Labial node is interpreted by the absence of the Dorsal node. Thus, the child may contrast the non-coronal segments /p~b~m/ with the coronal segments /t~d~n/ at this stage.

A second possible pathway of acquisition for a child who has initially acquired the Supralaryngeal node at Stage 1 is to elaborate the SV node before elaborating the Place node. This pathway would lead the child to contrast obstruents versus sonorants without reference to Place of articulation contrasts. Figure (8) presents Stage 1 with the elaboration of the SV node:

(8) *Stage 1: Addition of the SV node*

For a child who initially elaborates the SV node before the Place node, the place of articulation is the default bare Place node, or Coronal. However, contrasts between obstruents and sonorants are maintained. As the child contrasts sonorants and obstruents without reference to place of articulation, he or she may contrast /t~d~p~b/ versus /n~m/.

In the Second stage of elaboration, the possible pathways of acquisition increase exponentially from two (elaboration of the Place node or the SV node) to four. At this stage, a child who initially elaborated the Place node at Stage 1 could create further contrasts *within* the Place of articulation node by adding the feature Dorsal as a dependent of the peripheral node:

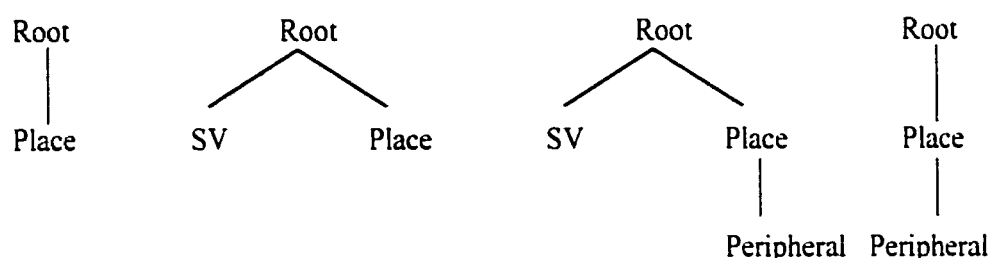
(9) *Stage 2: Addition of the Dorsal node*

At this stage, the child has fully elaborated the Place node without elaborating the SV node. Thus, contrasts of Place will be maintained without reference to distinctions

between sonorancy or obstruency so that the child will contrast the labial (peripheral) sounds /p~b~m/ versus the coronal segments /t~d~n/ versus the Dorsal sounds /k~g~ŋ/.

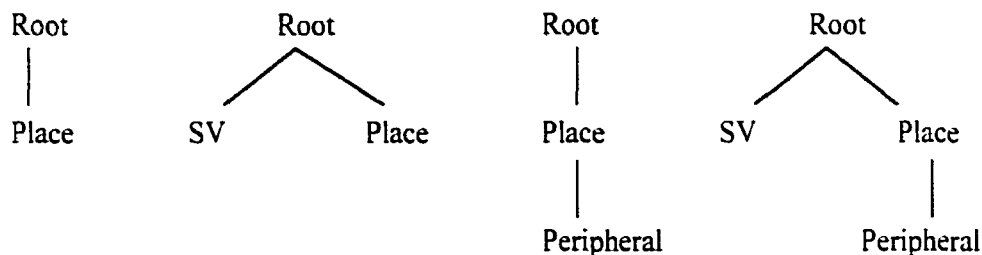
Alternatively, the child who had initially elaborated the Place node in Stage 1 could subsequently elaborate the SV node in Stage 2 to create a minimal obstruent versus sonorant contrast as seen in (10):

(10) *Stage 2: Addition of the SV node*



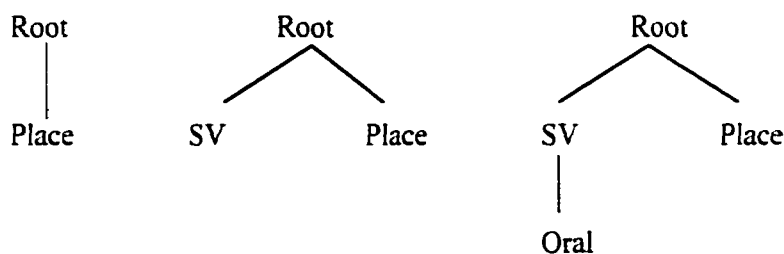
The child elaborating the SV node at this stage will be able to contrast sonorants versus obstruents within two place distinctions: coronals and non-coronals. Finer-grained distinctions within the Peripheral node have not yet been acquired. At this stage, the child will likely contrast the coronal obstruents /t~d/ versus the coronal sonorant /n/ and the non-coronal sonorant /m/ versus the non-coronal obstruent /ŋ/.

At Stage 2, children who initially elaborated the SV node in Stage 1 now have two possible pathways. One path is to elaborate the Place node to create contrasts between coronals and non-coronals for both sonorants and obstruents. Stage 2 representations with elaboration of Place by the feature peripheral are shown in (11):

(11) *Stage 2: Addition of the Peripheral node*

At the surface level, the representations shown in (11) yield the same possibilities for contrast as those presented in figure (10) for a child who has first elaborated the Place node in Stage 1 followed by the SV node in Stage 2. Both children would also be able to contrast the coronal obstruents /t~d/ versus the coronal sonorant /n/ versus the non-coronal sonorant /m/ versus the non-coronal obstruent /m/, even though the path of acquisition to this stage of the SV and Place nodes is reversed.

The fourth possible path of Stage 2 acquisition requires the addition of structure to the SV node by a child who had initially expanded the SV node in Stage 1. Further elaboration of the SV node would create a contrast *within* the class of sonorant sounds, that is, between nasal and nonnasal sonorants. Nasal sounds are represented by a bare SV node, while nonnasal sonorants require elaboration of the SV node by way of the Oral node. The representations are shown in (12):

(12) *Stage 2: Addition of the Oral node*

At this stage, the child has acquired quite complex contrasts within the SV node without reference to Place contrasts. In order to contrast coronal and non-coronal sounds, further elaboration of the Place node is required. A child who has elaborated the SV node in this manner distinguishes between the class of obstruents /p~t~k/ versus the nasal sonorants /m~n/ versus the oral sonorant /l/. Rice & Avery point out that this elaboration is less common than the other three at an early stage: however, their findings are based on the acquisition of English. Further cross-linguistic investigation may find that this path is preferred by learners of other languages⁴.

Models of structural elaboration can thus account for two types of variability: (i), variability resulting from a lack of phonological contrast due to the initial-state, non-expanded structure, and (ii), variability due to paths of acquisition inherent in the hierarchical structure of the feature geometry with its representations of dependence and constituency. The claim is that by incorporating variability into the model, arguments against deterministic models of language acquisition are mitigated. Moreover, the model of phonological acquisition as a process of structural elaboration makes the crucial assumption that children are creators of their own, individual yet constrained grammars and not merely mini-adults with flawed grammars.

Because the structure building hypothesis can account for both the variability and uniformity found in phonological acquisition, it is the model I adopt in this thesis. Having looked at Rice & Avery's arguments for segmental acquisition as a process of structure-building, we now turn to Brown's (1997: 1999) work which expands on and elaborates this approach by integrating the L1 phonological system in both L1 and L2 speech perception and production.

2.2 *(Brown (1997, 2000): The L1 phonological system in speech perception*

We now turn to work by Cynthia Brown and John Matthews (see Brown 1993, 1997, 2000 as well as Brown & Matthews, 1997) which provides the basis for the

⁴ Rice & Avery note that this path of elaboration may account for the early acquisition of laterals in Quiché (Pye, Ingram, & List, 1987).

research conducted in this thesis. Brown and Matthews take the claims of segmental acquisition as a process of structural elaboration a step further by incorporating the role of the L1 phonological system with the process of speech perception of both first and second language learners. In the following section I will discuss Brown's hypothesis as to the role of the language-specific phonological structure in speech perception for first language acquisition. This hypothesis provides the background for the model of L1 phonological interference in L2 phoneme acquisition that I discuss in Chapter Three.

2.2.1 Infant speech perception

In order for an infant to detect that two speech sounds are used contrastively, he or she must first be able to perceive the contrast. As Brown (2000: 14) points out, "proper development of the phonological system is dependent on properties of the speech perception mechanism. Given the fact that a child may be born into any language environment, it is imperative that he or she be equipped with adequate cognitive machinery to perceive (or, at the very least, be predisposed to perceive) the whole range of possible phonetic contrasts".

Seminal work by Kuhl, Werker and their respective colleagues (see for example Werker, 1981; Werker & Tees, 1983; Tees & Werker, 1984) has been highly influential in research on language-specific sound perception and the processes of phonological development in pre-linguistic infants, which is of interest to second language researchers as the construction of phonological representations is essential in both instances. Over the course of two decades, researchers have tested the ability of infants, children and adults in discriminating both native and non-native contrasts. Results show that discriminatory abilities vary with respect to age: infants as young as one month of age can effectively discriminate both native and non-native contrasts.

Werker and Tees (1984) found that 6-8 month old English infants were able to discriminate non-native contrasts between the Hindi alveolar [t] and retroflex [ɖ] as well as the Salish velar [k] and uvular [q], but that this ability was lost by 10-12 months of age by the English speaking infants *but not* by the Hindi and Salish speaking children for their respective languages. Using the head-turning method, researchers found that one-month old infants reacted to both the English phonemic contrasts as well as the non-

native contrasts like Hindi dental /t/ versus retroflex /ɖ/ and Salish glottalised velar /kʰ/ versus uvular /qʰ/, despite having no prior contact to either Hindi or Salish. Interestingly, however, at approximately 7 months of age infants begin to experience a decline in their abilities to discriminate non-native speech sounds (Werker & LaLonde, 1988). At around 10 months, the infants no longer reacted to the non-native Hindi or Salish speech sounds. Thus, it appears that the perceptual ability to discriminate non-native sounds *decreases* with exposure to the native language, while phonological ability to discriminate segmental contrast *improves* from no contrasts to only native contrasts.

However, this ability to distinguish non-native speech sounds is not lost for all contrasts, nor does the decline in perceptual discrimination occur at the same time for all contrasts. For example, Werker & Tees found that both children and adults were perceptually sensitive to Zulu clicks, but not to Hindi retroflex stops.

Brown (1993) points out that the temporally non-uniform decline in perceptual ability for non-native contrasts suggests that loss of perceptual sensitivity is gradual and systematic, and that an explanation integrating linguistic experience and speech perception is needed. To address this, Brown proposes that the decrease in ability to discriminate non-native sounds and the corresponding increase in ability to discriminate native language contrasts is linked to the same mechanism: the step-by-step elaboration of the segmental feature hierarchy in the L1 which imposes a template or filter through which non-native segments are perceived.

In the subsequent section we look at Brown's hypothesis as to the link between phonological development and acoustic discrimination and its implications for first language acquisition. Implications of the model for second language acquisition are discussed in the following chapter.

2.2.2 Brown: The link between phonological development and acoustic discrimination

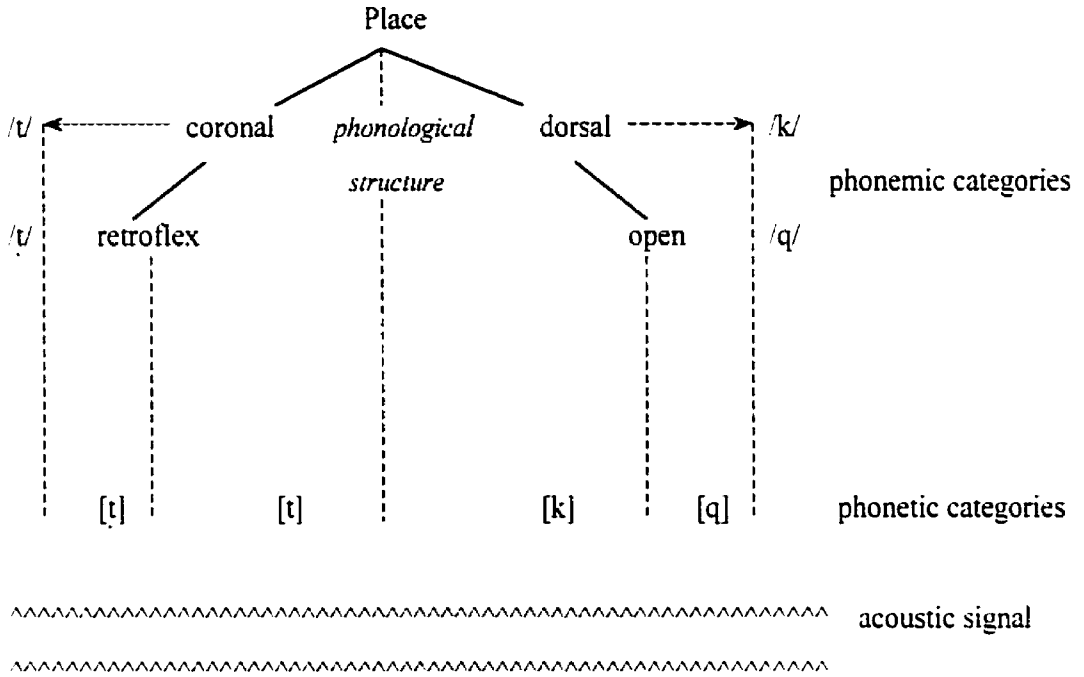
Based on the findings from infant speech perception showing that the decrease in perceptual ability to discriminate non-native sounds corresponds to an increase in phonological ability to discriminate native segmental contrasts, Brown proposes that the link between phonological development and acoustic discrimination can be accounted for

by the same mechanism. Essentially, the step-by-step elaboration of the hierarchical feature geometry in first language acquisition imposes a template or boundary on the perceptual system within which the language-specific phonemic categories are perceived. In this way, the detection of phonemic contrasts in the input language triggers the elaboration of the language-specific phonological hierarchy. This phonological structure thus acts as intermediary between the acoustic signal and the linguistic system by channeling the distinct acoustic signals into phonemic categories.

This filtering of the acoustic signal can help account for the well documented phenomenon of categorical speech perception, which is an all-or-none, or discontinuous, phenomenon (as opposed to the gradations found in continuous perception) in which sounds differing acoustically are mapped onto the same category. Thus, speakers of a particular language are better able to distinguish *between* members of different phonemic categories than *within* a single phonemic category (Pisoni, 1973; Repp, 1984). Supported by work done by Werker & Logan (1985) who found that some non-native contrasts may be perceived at either the auditory, phonetic or phonemic level depending on the length of the interval between stimuli, Brown's model proposes three different levels of processing: the auditory level, the phonetic level, and the phonemic level, suggesting that at some level of processing even non-native contrasts remain distinct to the hearer. In this way, perceptual reorganization accounts for loss of ability to discriminate non-native contrasts.

Figure 2.3 presents Brown's mechanism illustrating how the elaborated, language-specific phonological structure acts to channel distinct acoustic signals into phonemic categories:

Figure 2.3 Mediation of speech perception by phonological structure



The model shows how the sounds [t], [ɾ] and [k], [q] will be represented for a language with the corresponding phonemes. At the phonetic level, the signals for alveolar [t] and the retroflex [ɾ] remain distinct for all hearers whether the native language phonemically contrasts alveolar and retroflex stops or not. However, the fully specified phonological feature geometry funnels distinct acoustic signals into phonemic categories so that if a language exploits retroflex as a dependent of the coronal node (as in this language), then the phonetic sounds [t] and [ɾ] will be perceived as the phonemes /t/ and /ɾ/. Because there is featural structure under the Coronal node, the two distinct acoustic signals are filtered into separate phonemic categories. In contrast, if the Coronal node is unelaborated, then the two distinct acoustic signals would be shoe-horned into a single phonemic category of coronal sounds. The detection of the contrast in the ambient language triggers elaboration so that if there is no contrast, then structural elaboration is unnecessary. This is the case with English stops that contrast three places of articulation. Labial /p/, Coronal /t/ and Dorsal /k/ but do not have contrast *within* a single place of articulation.

In this chapter, I have outlined the framework for the acquisition of segmental phonology by first language learners, including arguments for the acquisition of speech segments as a process of structural elaboration. The motivation for the inclusion of a chapter on child language acquisition in a thesis concerned with adult second language acquisition is such that understanding the developmental process that children undergo when acquiring the phonological contrasts of a first language aids us in understanding the challenges adults face in acquiring the novel sound system of a second language. Also discussed was Brown's model linking phonological development and acoustic discrimination in first language acquisition, which influences how second language learners perceive novel contrasts in the target language. Thus, the assumptions made in this chapter provide a background to theoretical assumptions regarding second language learners' acquisition of new segmental contrasts to be discussed in Chapter Three upcoming.

CHAPTER THREE

SECOND LANGUAGE PHONOLOGICAL ACQUISITION

0.0 BACKGROUND

As we saw in Chapter Two, first language acquisition of segmental structure is motivated by the interaction between Universal Grammar and the process of structure elaboration in response to language specific phonemic contrasts detected in the input. The acquisition of the segmental sound system of a second language is different, however, as the second language learner comes to the acquisition task with a phonemic inventory already established in the L1 and must acquire the novel phonemic contrasts of the L2.

The existing sound system of the first language can affect both the perception and production of the second language. Research shows that L2 learners make perceptual reference to L1 phonetic categories in order to impose structure on L2 speech (see Abramson & Lisker (1970) for VOT; Bluhme (1969) for Australian perception of German vowels; Strange (1992) for Japanese perception of English approximants). Much of the work done by Flege (see for example Flege, 1987, 1988, 1990) suggests that L2 learners project their L1 phonetic categories whenever possible on the sounds of the L2 in a process of “equivalence classification”; this process of classification occurs even when there are detectable acoustic difference between the L1 and the L2.

As discussed in Chapter Two, in order to produce a sound segment of a language, the speaker presumably makes reference to some internal representation of that sound and then initiates the appropriate motor command according to production rules. However, in order to produce a relatively unknown or less familiar sound in the L2, the speaker must refer to a less well-formed or even inaccurate mental representation, and then enact the motor commands according to that (mis-)representation. Production is thus influenced at both the perceptual and the articulatory level. Research in the production of new sounds has shown that interference from the L1 phonemic sound system results in L2 speech patterns that are marked by non-native phonological patterns in the form of a foreign accent.

This chapter will discuss the acquisition of segmental phonology by second language learners. I begin Chapter Three with a theoretical overview of second language phonological acquisition. Section 2 discusses theories of language acquisition, including criteria of a good theory and discussion of early theories of second language phonological acquisition. Section 3 is an in-depth look at the model assumed in this thesis. Brown's (1997, 1998) model of L1 interference which links first language acquisition of phonological structures with difficulties encountered by the second language learner.

1.0 SECOND LANGUAGE PHONOLOGICAL ACQUISITION

Research in second language phonological acquisition requires integrating two areas of linguistic inquiry: the theory of acquisition and the theory of phonology. This integration of two distinct domains of linguistic inquiry has been aided by recent models of generative phonology, which are also inherently models of phonological acquisition. Thus, as we saw in Chapter Two, the Continuity Hypothesis assumes that children's construction of phonological representations are constrained and guided by principles of UG at each developmental stage. Although a child's developing grammar may differ at a particular stage from the target adult grammar, at no stage in the acquisition process will the grammar violate UG principles.

When we turn to the domain of second language phonological acquisition, however, a crucial difference between L1 and L2 acquisition quickly becomes apparent: the infant learning his or her first language appears to acquire the phonology effortlessly, while the phonology of an adult learning a second language is marked by non-native phonology in the form of an accent. Thus, two critical (if obvious) factors separate first and second language acquisition: (i), the first language learner comes to the task of language learning 'fresh', while the second language learner already has an established language which can both aid and hinder the acquisition of further languages, and (ii), many second language learners begin to learn the second language past the maturational point at which a native language is learned. This chapter is concerned primarily with the first point; that is, with the effects of prior linguistic knowledge on second language acquisition, specifically in how the existing L1 sound system impinges

on the acquisition of novel phonemic contrasts.

As we saw in Chapter Two, the logical problem of language acquisition led to the positing of an innate mental structure known as UG, which consists of a system of innate principles and parameters where principles are universal structures underlying all languages, and parameters are language specific characteristics triggered by the input (Chomsky, 1988). This innate system of linguistic universals greatly simplifies the task of first language acquisition in that the child comes to the task equipped with this set of universal structures, and thus does not need to overtly learn them. First language acquisition is thus seen as a process of the innate, UG-provided, principles and parameters interacting with the language-specific input the child is exposed to. However, as is made abundantly clear to anyone who has tried to learn a foreign language, adult second language acquisition is different than the relatively painless process of learning a first language. Given the difficulties of adult second language acquisition, the question arises as to whether or not UG is still accessible to second language learners, and, if so, to what extent. These issues are discussed in the following section.

1.1 Access to UG in Second Language Acquisition

In arguing that the rapid acquisition of a native language is due largely to the complex principles provided by UG interacting with the language specific input, we must then consider UG in relation to the slower, more painful acquisition of a second language. Is UG accessible to second language learners? To what extent? Although most research on the question of access to the principles and parameters of UG has been in syntactical acquisition, the same approach can be applied to phonological acquisition.

Research in phonological acquisition has shown that, in the area of stress assignment at least, learners may arrive at parameter settings that are not present in either the L1 or the L2 but that stress assignment is nonetheless UG-constrained in that learners do not arrive at 'wild' or 'impossible' grammars. However, in assuming the existence of UG we need to distinguish between questions of *what* and *how* in second language acquisition: (1), What is the nature of the L2 representation, and (2) How is it acquired (White, 1996).

There are two main positions with respect to UG availability: (i), the *No-Access*

approach, whose proponents argue that UG is not accessible to language learners after the critical period, roughly corresponding to the onset of puberty, has passed; and (ii), the *Access-Approach* which claims that learners have some access, whether full or partial, to UG. From these two positions, there are three possible hypotheses as to the role of UG in second language acquisition: No-Access; Partial-Access; and Full-Access.

The No-Access approach to UG in second language acquisition explicitly denies the accessibility of UG after the first language has been acquired. For example, both Bley-Vroman (1990) and Schacter (1989, 1996) (among others) argue that UG may guide language acquisition during a critical period before puberty, but that it is not available to an adult, or post-puberty, second language learner after this sensitive period has passed. That is, UG has a limited span which cannot be activated after the passing of the critical period and thus L1 and L2 acquisition are argued to be fundamentally different processes.

Supporters of the no-access to UG position argue that L2 learners often have difficulty in learning the grammar and sound systems of a new language, that acquisition is rarely perfect, and that the grammars of L2 learners would be more uniform if they had access to UG. Problems in the acquisition of a second language are explained by assuming that UG is no longer available to the second language learner so that any learning that does take place comes as a result of inductive learning strategies such as memorization and problem solving strategies, and cannot be attributed to an innate language module such as UG. Moreover, similarities between child L1 learners and adult L2 learners can be attributed to the fact that there are more similarities than differences among languages so that adult L2 learners are familiar with the characteristics of language and this aids in acquiring a new language.

The No-Access approach is problematic in that it fails to account for what second language learners know about a language. By relying on inductive learning strategies such as memorization in the L2 acquisition process, the No-Access hypothesis is merely describing the tools of language acquisition, and not the properties of the end-state grammar as it is represented in the brain of the learner. In contrast, the mental properties of the end-state grammar are precisely what a theory of UG provides us with. Also, the no-access approach to UG would predict that we would see wild, or unconstrained, grammars in L2 acquisition. However, research has shown that while nonnative speakers

develop interlanguage grammars that differ from the target grammar, these interlanguage grammars are nonetheless constrained by UG; that is, they are not wild grammars (see Broselow & Finer, 1991; White, 1992). On this basis, in this thesis I reject the No-Access hypothesis in favor of the Partial-Access Hypothesis, where UG is argued to be accessible to the second language learner in the form of the principles and parameters instantiated in the L1.

Proponents of the role of UG in second language acquisition argue that UG is still accessible to the adult second language learner. Arguments supporting the availability of UG (whether partial or full) to second language acquisition focus on the logical problem of second language acquisition: since L2 learners acquire complex mental representations that are not identical to those of the target language, how can we account for this in light of the input? (White, 1985, 1989). Under the access to UG framework, language learners are seen as constructing their own grammars which are not identical to the target language, but are nonetheless subject to and constrained by principles of UG. These individual grammars, termed *interlanguages* by Selinker (1972), may or may not have some properties in common with the target language, but they could be possible natural languages since they are constrained by UG. Thus, first and second language acquirers have different competencies arrived at by the same means of acquisition. Arguments for adult access to UG focus on interlanguage grammars of second language learners to see if they are constrained by principles of UG.

Supporters of the access to UG approach point to current research which has challenged the existence of the critical period (Flynn & Martohardjono, 1994), and has shown that there may be an underlying commonality to both child L1 acquisition and adult L2 acquisition. Moreover, it has been shown that linguistic experience underdetermines linguistic knowledge for the L2 learner as much as for the L1 (Cook, 1988; Gregg, 1996; White, 1989, 1996). That is, there exists also a logical problem for L2 acquisition, so that we should expect the L2 acquisition process to be constrained by innate principles similar to those of L1. Although researchers may largely agree that there is a logical problem in L2 as in L1 acquisition, there is a difference of opinion as to whether the solution to the problem is the same: that is, researchers differ as to whether UG is unavailable, fully available, or only partially available. Thus, access to UG

proponents can be divided into two groups: those supporting the strong version claiming full access to UG by L2 learners; and those supporting the weak version which argues for partial access to UG in the form of parameters instantiated in the L1.

The *Full-Access* hypothesis claims that all principles and parameter values available to the child via UG are still fully accessible to the adult L2 learner. Under this claim, the differences in patterns of acquisition between first and second language learners can be explained in other ways than by positing a lack of access to UG. This position is espoused in the work of Finer & Broselow (1986); Flynn (1987, 1991, 1993); and Martohardjono (1991, 1993) among others.

However, the Full-Access approach is problematic in that the ultimate attainment of most adult L2 learners is not equal to that of first language learners. Also, in what has been called cross-linguistic transfer, the L1 can either facilitate or hinder the acquisition of a new language, depending on the similarities and differences between the L1 and the L2.

The *Partial-Access* hypothesis claims that UG is accessible to the second language learner only partially: UG is thought to be accessible to the L2 learner only in the form of the principles and parameters of the L1 that are present in the L2. That is, UG does not constrain adult L2 hypotheses as it does for the child L1 learner; rather the grammar constructed by the L2 learner is mediated by the parameters set in the native language. Ultimate attainment is impossible in those cases where L1 and L2 parametric values are mismatched. As will be made explicit shortly, the model of L2 phonological acquisition adopted in this thesis assumes partial access to UG in claiming that the phonological systems of L2 learners are constrained by UG and follow a developmental path mediated by the phonological feature structure of the native language.

In the next section we discuss the requirements of a good theory of language acquisition before turning to Brown's (1997, 1998) model of L1 interference in L2 acquisition. Outlining the characteristics of a good theory of language acquisition will allow us to see why Brown's model is able to account for the patterns in the acquisition of a second language sound system.

2.0 MODELS OF SECOND LANGUAGE FEATURE ACQUISITION

2.1 *Characteristics of a good model of second language acquisition*

While the assumption of the existence of a language learning module such as UG helps to account for the logical problems of both first and second language acquisition, we still require a theory to account for *how* language acquisition happens. Gregg (1996) proposes three essential criteria for a good theory of second language acquisition:

1. The Theoretical Framework Criterion:

The theory must be constructed within the framework of a unified general theory. We don't want to have radically different property theories each accounting for knowledge at a different development stage.

2. The Sequence Criterion:

The order of acquisition must be explicable. That is, we must be able to account for why X occurs before Y and not vice versa in a developmental sequence.

3. The Mechanism Criterion:

There must be a detailed specification of the acquisition mechanism.

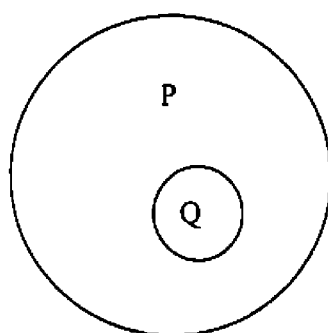
Let's go over the criteria in greater detail. The *Theoretical Framework* criterion requires a general learning theory for the acquisition of language, and not various unrelated sub-theories accounting for each developmental step. As Long (1993) points out, there are over 60 theories of language acquisition including theories of variation, production, interaction and so on. Given the various subdomains of linguistics, we cannot easily assume one single theory to account for all areas, however, a set of complementary theories to account for the language learners' competence is required. Current linguistic theory has proposed models integrating theoretical linguistics and second language acquisition in which the learners are seen as constructing individual, UG constrained grammars, or interlanguages.

The *Sequence Criterion* requires that the developmental order of structures must be accounted for in any theory of language acquisition. A good theory must be able to

explain why a particular order of acquisition reflects the learners' internal grammar since interlanguages are argued to be constrained by principles of UG. As seen in Chapter 2, phonological models of phoneme acquisition assume a hierarchical expansion of the Feature Geometry based on structural relationships of constituency and dependency to account for the order of acquisition of segments.

Finally, the *Mechanism Criterion* requires that there be some learning mechanism operating on linguistic input. In a modular framework, the most commonly held learning principle is the Subset Principle which constrains hypotheses on the basis of markedness assumptions. Because research in first language acquisition suggests that children either do not receive or do not notice negative evidence (e.g. Brown & Hanlon, 1970), it is argued that language must thus be acquired mainly on the basis of positive evidence, and that the learner must be somehow prevented from arriving at over-inclusive grammars requiring negative evidence. In a modular generative framework, possible hypotheses are constrained by UG and the order of acquisition is constrained by markedness assumptions based on the Subset Principle (Berwick, 1985; Manzini & Wexler, 1987; Wexler & Manzini, 1987). To illustrate the Subset Principle, consider 2 grammars which are in a Subset/ Superset relationship where both grammars generate the same subset of sentences and one of the grammars generates additional ones. The Subset Condition is shown in (1):

(1)



The grammar that generates the set of P sentences also generates the Q set of sentences so that Q is a subset of P but not vice versa. The learnability problem occurs because the set of Q sentences are compatible with two grammars: the grammar that generates the Q

sentences as well as the grammar generating the P set. If a language learner is learning a language that contains only the set of Q sentences, the appropriate hypothesis is the grammar which generates Q but not P, since hypothesizing the P grammar will result in overgeneralizations that cannot be disconfirmed on the basis of positive evidence. The grammar generating the subset is the unmarked, default option while the grammar generating the superset is marked and thus hypothesised only on the basis of positive evidence.

It has been argued that UG generates languages in a subset/superset relationship meeting the Subset Condition (Berwick, 1985; Wexler & Manzini, 1987; White, 1989). The learnability problem is overcome by the Subset Principle which can be seen as a constraining mechanism: given input which can be generated by either of two grammars meeting the Subset Condition, the learner should adopt the most restrictive grammar consistent with the input. The superset grammar should only be hypothesised on the basis of positive evidence.

Research in the acquisition of phonological parameters has been concentrated in the area of stress assignment, investigating the acquisition of English word stress by speakers of languages with different parametrical settings to see which aspects of the input might act as cues for resetting the metrical parameters (see Archibald, 1992, 1993; Dresher & Kaye, 1990; Pater, 1993).

Having discussed the minimal criteria for a good theory of language acquisition, we now turn to some proposed theoretical models of second language acquisition. Discussing the strengths and shortcomings of earlier models will give us insight into why Brown's model of L1 interference (adopted in this thesis) provides an improvement over previous models.

2.2 *Theoretical Models of Second Language Acquisition*

Over the past three decades as the fields of second language production and perception have expanded, researchers have proposed various theoretical models in an attempt to explain how the novel sounds of a second language are mapped onto, or perceived in terms of, the native phonemic inventory. These models focus on linguistic explanations for non-native speech patterns in second language speech; however, it must

be noted that other factors also affect a learner's progress and ultimate achievement in second language speech. These factors may include such things as maturational factors; social factors including social acceptance (Brennan & Brennan, 1981) or social distance; and individual factors including motivation (Purcell & Suter, 1980), gender (Weiss, 1970), personality (Heyde, 1979), and oral and auditory capabilities (Locke, 1968, 1969).

As background to the recent model of phonological interference assumed in this thesis, I will describe three earlier models of second language phoneme acquisition which revolve around questions of how novel sounds are mapped onto existing L1 categories: the Perceptual Assimilation Model (PAM) (Best, 1993, 1994); the Speech Learning Model (Flege, 1991, 1995); and the Feature-Competition Model (FCM) (Hancin-Bhatt 1994). In terms of a language acquisition theory, each of these models has its own strengths and weaknesses in characterising the nature of segmental acquisition, and I will go over each in turn to see how Brown's model offers an improvement to previous theories.

2.2.1 Perceptual Assimilation Model (PAM) (Best, 1993, 1994)

Best's Perceptual Assimilation Model (1993, 1994) argues that non-native speech sounds are assimilated to L1 phonemic categories on the basis of their articulatory similarity. Articulatory similarity is defined as the spatial proximity of the location of the constriction of the active articulator between the native language and the target language. Based on articulatory similarities, if an L2 sound can assimilate to an existing L1 phoneme then the learner will be able to perceive it.

While this model addresses the relationship between speech perception and the native phonological sound system, it fails the second and third criteria proposed by Gregg for a model of language acquisition: namely, it does not identify a developmental sequence in which the contrasts may be acquired nor does it provide a mechanism or criteria to determine exactly how non-native contrasts assimilate and are acquired.

2.2.2 Speech Learning Model (Flege, 1991, 1995)

Flege's Speech Learning Model distinguishes between 'new' and 'similar' sounds, where the term "sounds" refers to the phonetic rather than the phonemic unit.

Because the term “new sound” may be potentially problematic since the difference between an L1 and an L2 sound may be a minor phonetic variation, Flege points out the importance of explicit and objective criteria for deciding what is a ‘new’ sound and what is a ‘similar’ sound for the second language learner.

New sounds in a second language are those that are not identified with any L1 sounds, while similar sounds are perceived to be the same as a particular L1 sound. Flege argues that the “phonetic distance” of the contrastive sounds predicts which novel segments will be acquired: “new” L2 sounds will be acquirable since they do not correspond to a sound in the L1 while “similar” sounds are perceived, and thus produced, incorrectly as the corresponding L1 sound rather than as the novel L2 sound. Identical phonemes are not problematic for acquisition. It is the process of ‘equivalence classification’ (Weinreich, 1953) that hinders the creation of new phonetic categories for similar sounds. This process of equivalence classification accounts for why English speakers perceive clicks (new sounds) but not retroflex segments (similar sounds).

Flege’s Speech Learning Model is problematic, however, as it relies on the phonetic unit rather than the phoneme. The question thus arises of how the phonetic unit maps to the phonemic level. As Leather and James (1996: 289) note: “A speech signal or phonetic interpretation of interlingual identifications must make a connection with a phonological interpretation of that part of an L2 speech learner’s “mental grammar” as much as the latter must connect with the former as a specification of part of a learner’s speech processing arsenal.”

2.2.3 Feature Competition Model (FCM) (Hancin-Bhatt, 1994)

Hancin-Bhatt’s (1994) Feature Competition Model (FCM) offers a connectionist (that is, nonmodular) perspective to the modular theories of speech perception and phonological acquisition we have been looking at. The Feature Competition Model is able to predict which features or segments will transfer from the L1 to the L2 on the basis of language specific feature prominences, calculated by a specific metric. The frequency of occurrence of a particular feature dictates its prominence in that language so that features do not have discrete values but rather vary in prominence according to language specific phonemic inventories. Features that occur more frequently (and are thus more

prominent) in a language will have a greater influence on the perception of new sounds in the target language; thus, the feature prominence of the native language phonemic inventory guides how novel sounds are mapped onto existing L1 categories.

This theory fulfills all three parameters for a good model of language acquisition: it is (i), formulated within an existing Theoretical Framework; (ii), has an explicable order of acquisition; and (iii), specifies an acquisition mechanism. However, the FCM does not address the interrelationship between the native language phonological system and perception of phonemic contrasts. Brown's model, to be discussed in Section 3, directly addresses the nature of phonological transfer from the L1 to the L2 by linking speech perception to the L1 phonological system.

3.0 BROWN'S MODEL OF L1 INTERFERENCE (1993, 1997, 2000)

In Chapter Two we discussed Brown's model linking phonological development and acoustic discrimination in the native language. What are the implications of this model for the acquisition of a second language sound system? The basic assumption of Brown's model is that it is the language-specific hierarchical structure of the native language that constrains or delimits which contrasts the learner will be able to accurately perceive, and hence acquire. Acquisition of second language phonemic contrasts is thus dependent on accurate perception of those contrasts, guided by the L1 feature geometry. Because the L1 feature structure is constrained by UG, so too will the interphonology of the L2 learner be constrained, based on the hierarchical structure of the feature geometry.

Brown's model of L1 interference thus integrates insights from speech perception and first language acquisition and applies them to second language acquisition in addressing how the L1 phonological system affects L2 speech perception. Her experimental goal is to provide evidence for UG accessibility and the role of the L1 in second language phonological acquisition. Three questions are addressed:

1. If UG is indeed active, how can we account for the failure of second language learners to acquire the phonological properties of the L2?
2. Which aspects of the first language interfere with L2 phonological acquisition?

3. Can L2 learners acquire new phonological representations under the right conditions?

To address these questions, Brown develops a model of second language phonological acquisition based on L1 interference whereby the intake to the language acquisition device is determined by the phonological structure of the first language. A crucial distinction is made between intake and input (see Corder (1967) for the distinction; also Carroll (1999) for an alternative position) where *input* refers to the language surrounding the learner while *intake* is the part of the input to which the learner is sensitive, or, in other words, the actual input to the language acquisition device.

The claim of the model is that the monotonic acquisition of the language specific feature geometry of the native language restricts sensitivity to, or intake of, non-native speech sounds by children acquiring a language. Furthermore, it is this existing phonological structure that then constrains which non-native phonological contrasts a L2 learner will be sensitive to, and hence able to acquire. Put another way, the L1 phonological system or feature hierarchy constrains the boundaries within which the non-native phonemes are defined. Thus, the phonological structure of the L1 acts as a sort of template defining the categories within which the L2 sounds are perceived. However, it is not the segments themselves that determine perception, rather, it is the segmental sub-units, or features, of segments that play an essential role in the acquisition of L2. Whether or not a L2 learner can perceive the non-native phonological contrasts is dependent on the specific feature structure of the L1.

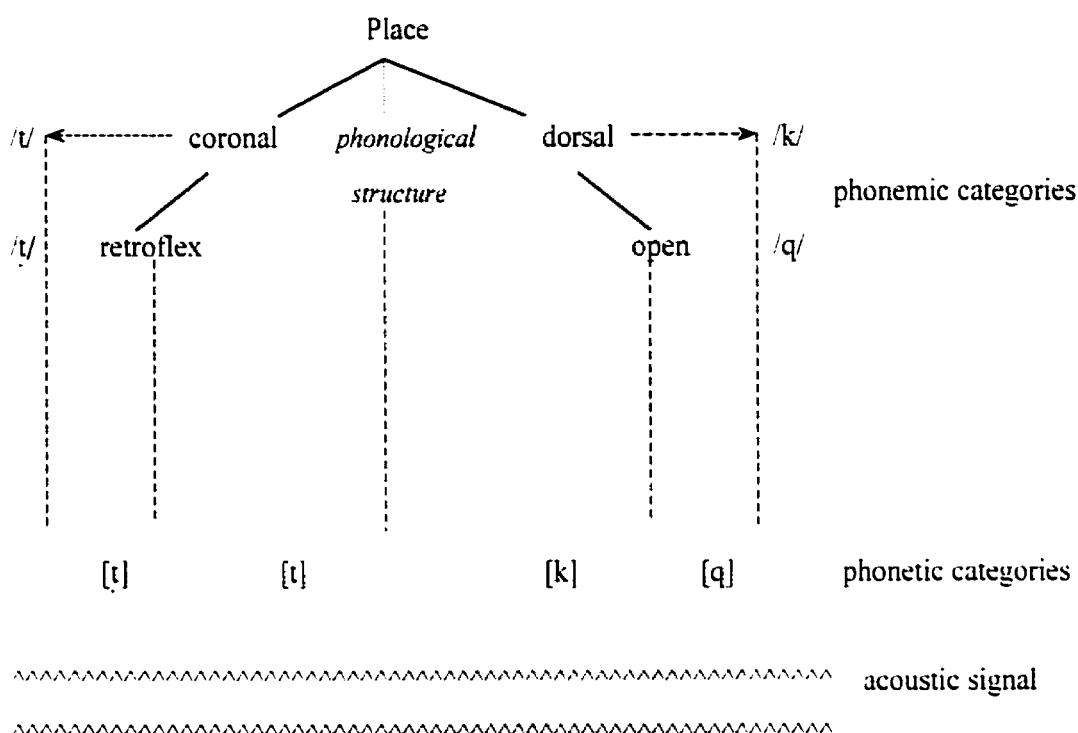
As the perception of a novel contrast is essential to establishing the structural representations phonologically characterising the two segments, the L1 set of phonological contrasts may actually impede accurate perception of the new phonemes, thereby inhibiting the acquisition of new structural representations. Specifically, if a non-native segment is characterised by a particular feature, then the learner will be able to perceive the non-native contrasts if he or she manipulates the feature elsewhere in the L1 grammar. That is, speakers of a given L1 can perceive non-native contrasts distinguished by a feature present (elsewhere) in their L1 grammar. As Brown notes: "despite a lack of acoustic, phonetic or phonemic experience with a *particular* non-native contrast, a

speaker's experience perceiving native phonemic contrasts along an acoustic dimension defined by a given underlying feature (for example, voicing) permits him or her to accurately discriminate *any* non-native contrast that differs along that same dimension." (22). Perception of a new phonemic contrast is facilitated by the presence of the distinguishing feature elsewhere in the learner's L1 inventory. However, if a particular feature is lacking in the representation of any phonemic contrasts in the learner's native feature geometry, then perception of the novel phonemic contrast should not occur. Moreover, because perception is crucial to establishing new phonological representation, the learner will not be able to construct a new representation for that segment.

3.1 *The phonological filter*

As we saw in Chapter 2 and Section 3.0 above. Brown argues that it is the featural, rather than the segmental level, which determines the perception of a phonemic contrast. Thus, the mapping of the acoustic signal into phonemic categories is guided by language specific featural makeup. In the acquisition of the native language, the feature structure is used to funnel the distinct phonetic variations into individual phonemic categories so that perception of non-native contrasts gradually declines as the novel sounds are interpreted as phonetic variants of existing categories. The model representing the phonological filter is repeated here as Figure 3.1:

Figure 3.1 Mediation of speech perception by phonological structure



Recall that the model illustrates how the language-specific phonological structure channels the distinct acoustic signal into phonemic categories, decreasing sensitivity to non-native contrasts in the process. In the acquisition of second language phonemic system, the same process of filtering occurs. At the phonetic level, the signals for the phonetic sounds remain distinct for all hearers, regardless of the native language. This claim for universal discrimination at the phonetic level is supported Werker & Logan (1985) who found that the speech signal is processed at three distinct levels (auditory, phonetic or phonemic) depending on the length of the interval between stimuli. Specifically, Werker & Logan showed that under certain conditions English speakers could discriminate the Hindi alveolar /t/ versus retroflex /ṭ/ contrast at greater than chance levels.

Although the acoustic signal is initially divided into distinct phonetic sounds,

once at the phonemic level the feature geometry funnels the distinct signals into phonemic categories according to the language-specific feature geometry. This channeling of various phones into discrete phonemic categories is thought to aid in the rapid processing and comprehension of speech. Given that the primary goal upon hearing a speech signal is to access higher level information at the semantic or lexical level, it is desirable that as little effort as possible is given to deciphering the lower level speech signal. Categorical perception allows us to filter out phonetic variation provided by coarticulation effects or inter-speaker variability to redirect attention to other tasks. However, this filtering of the acoustic signal to aid processing can negatively influence the perception of a non-native language. As Brown (2000: 19) states: "variation in the acoustic signal which is filtered out by the native phonological system (i.e., is treated as *intra*-category variation) may, in fact, contribute to differences in meaning in the foreign language (i.e., actually constitute *inter*-category variation)." Thus, the representations developed in the course of first language acquisition affect how the non-native phonemes are perceived.

3.2 *Implications for L2 acquisition*

Brown's model makes different predictions for the acquisition of novel phonemic contrasts depending on whether one, both, or neither member of a phonemic contrast is present in the native language. There are three ways in which a non-native segmental contrast can correspond to the L1 inventory.

The first type of L2 phonemic opposition occurs when both members of a non-native contrast correspond to segments in the native language. The two segments may either match exactly the corresponding segments in the L1: for example, the voiceless bilabial stop [p] may contrast with the voiced counterpart [b] in both the L1 and the L2. It may also be the case that the contrastive pair in the L2 is similar to but not exactly the same as the contrasting phonemes in the L1: for example, for an English speaker learning Salish, the English /t/ ~ /k/ contrast would correspond to the Salish glottalised /t̚/ versus glottalised /k̚/. In this case, the native English speaker would categorise the contrastive Salish glottalised stops /t̚/ ~ /k̚/ as members of the English contrastive phonemes /t/ ~ /k/. In any case, a between-segment contrast is maintained acoustically

even though the non-native segments are likely represented by the existing L1 structural representations and no new structure is being added.

The second type of non-native contrast encountered in L2 acquisition occurs when neither phoneme of a non-native contrast is present in the L1 phonemic inventory. Due to constraints on possible phonological contrasts, it is quite rare to find a phonemic opposition where neither member is present in the L1; however, one such opposition can be seen in the acquisition of Zulu clicks by native English speakers. Because clicks do not correspond to any segment in the English phonemic inventory, English speakers will not have a featural structure to represent the segments and so the perception of the contrast will not be inhibited by existing L1 structure. Two predictions based on UG access are made as to whether the novel phonological representation can be acquired: if UG is still active in L2 acquisition, then the learner will be able to build a new featural representation for the segmental contrast on the basis of perception of the novel contrast. If, however, UG is not active then the learner will not be able to construct the new phonological representation despite perceiving the contrast.

The third type of L2 phonemic opposition is where one of the members of a non-native contrast is present in the native phonemic inventory, but the other segment is not. Leather & James (1996: 276) note that it is probable that "the adult's L2 phonetic learning task is harder for a sound classified as equivalent to one found in L1 than for one for which a phonetic category must be constructed from scratch -- because the influence of the L1 category may cause learners to develop inaccurate perceptual targets for the L2".

With respect to this third type of phonemic opposition, Brown's model makes two predictions as to the likelihood of acquisition of the non-native phoneme depending on the feature structure of the native language. As the L1 feature geometry funnels acoustic signals into phonemic categories, perception of the novel L2 contrast is determined by the presence or absence of the relevant feature in the first language. If the native language feature geometry does not manipulate the distinguishing feature, then perception of the L2 contrast will be inhibited by the L1 phonological structure. It must be emphasised that the feature need not be used in representation of the same segment; it can be used in the representation of another segment and will still be present as a building block for the

representation of a new segment. However, if the L1 feature geometry lacks absolutely the relevant feature in the structure of any segment, then perception of the novel segment will be blocked. Thus the possibility of acquisition is dependent on the featural, not the segmental, level of the L1. If the feature does exist in the native feature geometry, then it can be used as a building block in the perception of other segments.

This third type of phonemic opposition is most relevant to Brown's work, as predictions are made on the basis of present or absent features. In the next section we will look at examples of this type of opposition and, as a precursor to the acquisition of the Czech alveolar versus palatal stop contrast discussed in Chapter Five of this thesis, we will look at Brown's (1998) experimental work supporting the hypothesis that the particular feature geometry of a native language guides the acquisition of an L2 segmental contrast.

3.3 *The experimental data (Brown, 1998)*

3.3.1 *Experiment 1: Japanese and Chinese speakers' acquisition of the English /l/ versus /r/ contrast*

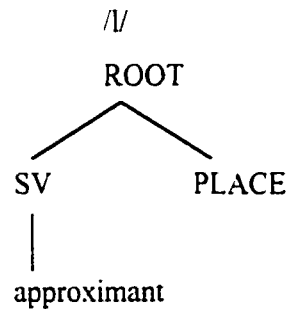
In order to provide experimental evidence for the claims of the phonological filter model, Brown (1998) investigated the acquisition of the English lateral approximant /l/ versus the central approximant /r/ by native speakers of Japanese and Mandarin Chinese. Speakers of these two languages were chosen because, crucially, neither Japanese nor Chinese phonemically contrasts /l/ and /r/ but both languages have a single phoneme that corresponds to one of the English liquids.

Japanese has a single liquid phoneme, described by Maddieson (1984) as a flap /ɾ/: the phonetic realization of this segment as an [l] or an [r] varies freely (International Phonetic Association, 1979; Vance, 1987). Because the two phonetic variants are in free variation in Japanese, they require only a single underlying representation (unlike the English /l/ versus /r/ contrast). The Mandarin Chinese inventory contains the lateral approximant phoneme /l/, whose phonetic variants do not vary between /l/ and /r/¹.

¹ This claim requires some comment. The Mandarin Chinese inventory contains a segment which is often transcribed in romanized script as "r". However, linguists (see Maddieson, 1984) classify this "r" as a voiced retroflex fricative, /ʐ/ and not as a retroflex sonorant.

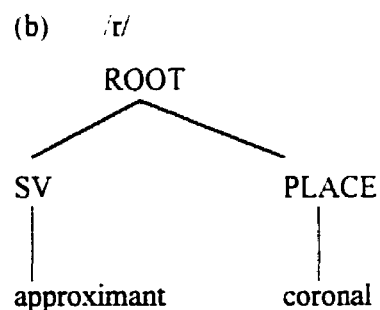
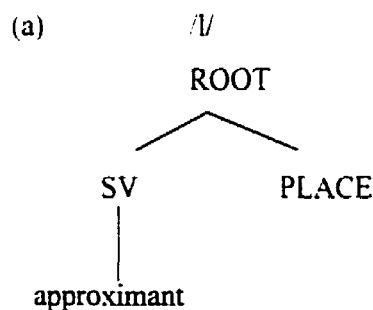
Because Japanese and Mandarin each have a single approximant, they do not require complex structural elaboration to distinguish between two approximants. The underlying representation of both the Japanese /ɾ/ and the Chinese lateral /l/ is shown in (2):

(2) Japanese approximant /ɾ/ and Mandarin approximant /l/



Because there is no phonemic contrast between approximants in either Japanese or Mandarin, all approximants are perceived as approximant segments with no finer grained distinction between them. However, because English does contrast a lateral approximant and a central approximant it requires further structural elaboration to distinguish between the two segments. Brown argues that the feature distinguishing the lateral approximant /l/ from the central approximant /ɾ/ is the feature [coronal] under the PLACE node in the representation of /ɾ/. Structures for the English approximants are given in (3):

(3) English structures for /l/ and /ɾ/:



Thus, native speakers of Japanese and Mandarin learning English will have to acquire the more elaborate structural representation for approximants in order to accurately produce the phonemic contrast between English /l/ and /r/. But in order to acquire the representations, they must first be able to perceive the contrast between the two segments. Brown's model claims that if a language manipulates a feature elsewhere in the inventory (independently of the segment at hand), then the learner should be able to perceive the contrast between the two non-native segments since they will have the building blocks required for the representation of that segment. Brown found that Mandarin but not Japanese requires the feature [coronal] in the representation of other phonemes: specifically, Mandarin requires the feature [coronal] as a dependent of the Place node in order to distinguish between two coronal sibilants, the alveolar /s/ and the retroflex /ʂ/. The phonemic inventory of Japanese does not require the feature [coronal] anywhere. The inventories of Japanese and Mandarin are in (4):

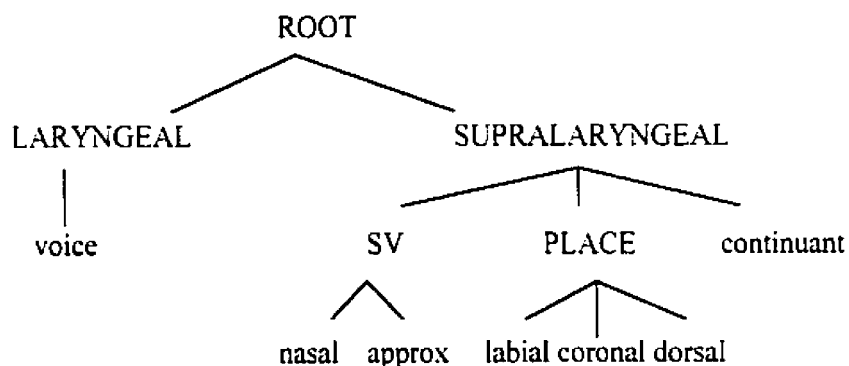
(4)	<u>Japanese inventory</u>			<u>Mandarin Chinese inventory</u>			
	p	t	k	p	t		k
	b	d	g	p ^h	t ^h		k ^h
		tʃ			ts	ʈʂ	
		dʒ		f	s	ʂ	h
		s	h			z	
		z		m	n		
	m	n			l		
		r		w	j		
	w	j					

As it is the composition of the language-specific feature geometry, and not the individual segments per se, that determines whether or not a novel L2 segmental contrast will be acquired, I present the fully specified adult feature geometries for Mandarin Chinese and

Japanese in (5) and (6), respectively:

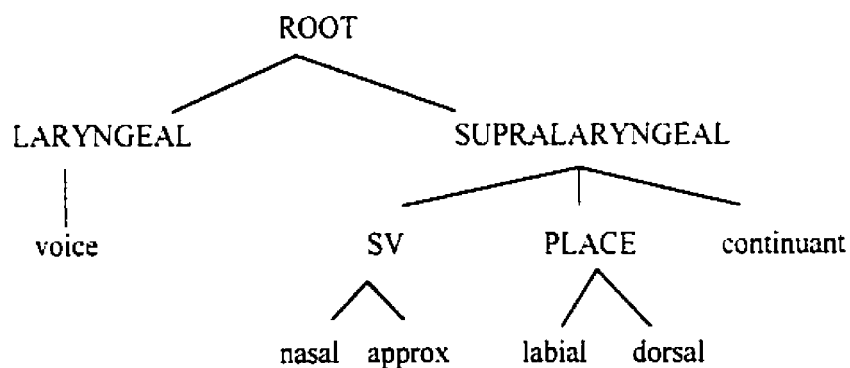
(5)

Chinese Feature Geometry



(6)

Japanese Feature Geometry



On the basis of these language-specific feature geometries, two predictions are made with respect to the acquisition of the English /l/ versus /r/ contrasts by native speakers of these languages. Even though neither Japanese nor Mandarin contrasts /l/ vs /r/ phonemically, (i). Mandarin speakers should be able to distinguish English /l/ and /r/ since they have the feature [coronal] elsewhere in inventory and so the acoustic signals for /l/ and /r/ will be mapped onto **separate** phonemic categories; and (ii). Japanese speakers won't be able to distinguish English /l/ and /r/ since they lack the feature [coronal] in their featural

inventory. For Japanese speakers, the acoustic signals for /l/ and /r/ will be shoe-horned into a **single** phonemic category of approximant, in the way that phonetic variants of a single phoneme are filtered into a single phonemic category.

To test these two predictions, Brown ran two tasks on two experimental groups: 10 native Chinese speakers learning English as a second language and 10 native Japanese speakers learning English as a second language, as well as 10 monolingual English speakers as a control group.

The first task was an AX discrimination task where subjects are presented with recorded cues of minimal pairs distinguished by the phoneme in question and asked to indicate whether the words are the same or different. This tests the subjects' ability to *acoustically* discriminate English /l/ from /r/.

The second task was a Forced Choice Picture Selection task where subjects are presented with two pictures as well as a verbal cue corresponding to one of the pictures: their task is to indicate which of the pictures the verbal cue is naming. This task tests whether or not the subjects who were acoustically able to discriminate the two segments can identify tokens of a particular phoneme, thereby indicating that they had acquired the phonological structure necessary to discriminate the two segments.

Brown's findings supported her hypothesis that it is the presence or absence of a distinguishing feature in a language learners native phonemic inventory that determines sensitivity to a novel contrast: although neither Japanese nor Mandarin Chinese contrasts the lateral approximant /l/ with the central approximant /r/, the two groups differed greatly in their ability to discriminate these contrasts. A factorial ANOVA showed highly significant differences between the two groups for both the onset condition ($F(2,27) = 171.025, p = .0001$) and the cluster condition ($F(2,27) = 71.381, p = .0001$). Post hoc Scheffe tests ($p < .05$) indicated that there was no significant difference between the Chinese and control groups, while the Japanese group performed significantly worse than both the Chinese and the control groups on these two conditions². Mean results for the AX Discrimination task are presented in Table 3.1:

² Carroll (1999) takes issue with Brown's claim that the Japanese subjects were unable to acoustically contrast English /l/ vs /r/. Carroll notes that the Japanese subjects were extremely accurate on the task when the segments occurred in coda position; she takes this as evidence that they seemed to be encoding the acoustic distinction in some way but are not employing the information for lexical selection.

Table 3.1: Mean Results of Discrimination Task: English /l/ vs /r/

	<i>Percent Correct by Position</i>		
<i>Language</i>	Onset	Cluster	Coda
Japanese	31.3%	38.1%	99.3%
Chinese	97.5%	89.2%	98.3%

Results from the Picture Selection task support the findings of the Discrimination task: the Japanese speakers performed significantly worse than the Chinese speakers. A factorial ANOVA revealed highly significant differences between the two groups for both the onset condition ($F(2,27) = 43.74, p = .0001$) and the cluster condition ($F(2,27) = 41.524, p = .0001$). Post hoc Scheffe tests ($p < .05$) revealed that the Japanese group differed significantly from both the Chinese and the control group on these two conditions, while the Chinese and control groups did not differ significantly from each other. Mean results are presented in Table 3.2:

Table 3.2 Mean Results of Picture Selection task: English /l/ vs /r/

	<i>Percent Correct by Position</i>		
<i>Language</i>	Onset	Cluster	Coda
Japanese	59.2%	50.8%	92.5%
Chinese	100%	90.8%	95.8%

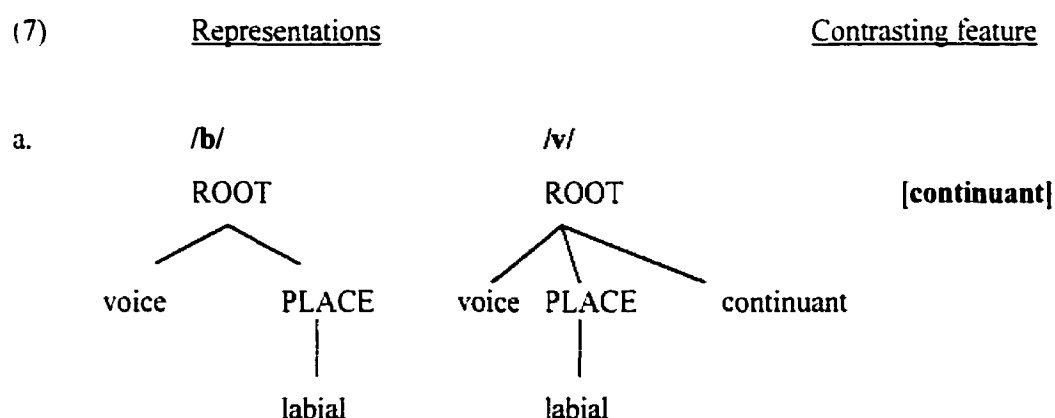
In sum: the results of these two experiments indicate that Japanese speakers are unable to contrast English /l/ versus /r/ either acoustically or phonologically while Mandarin Chinese speakers perceive this contrast with native-like accuracy. These results support Brown's hypothesis that it is the L1 feature inventory which determines whether or not a novel segmental contrast will be perceived by L2 learners. Brown

further tests this hypothesis by comparing Japanese speakers' acquisition of other non-native English contrasts differing in a single feature. We look at the results of this experiment in the next section.

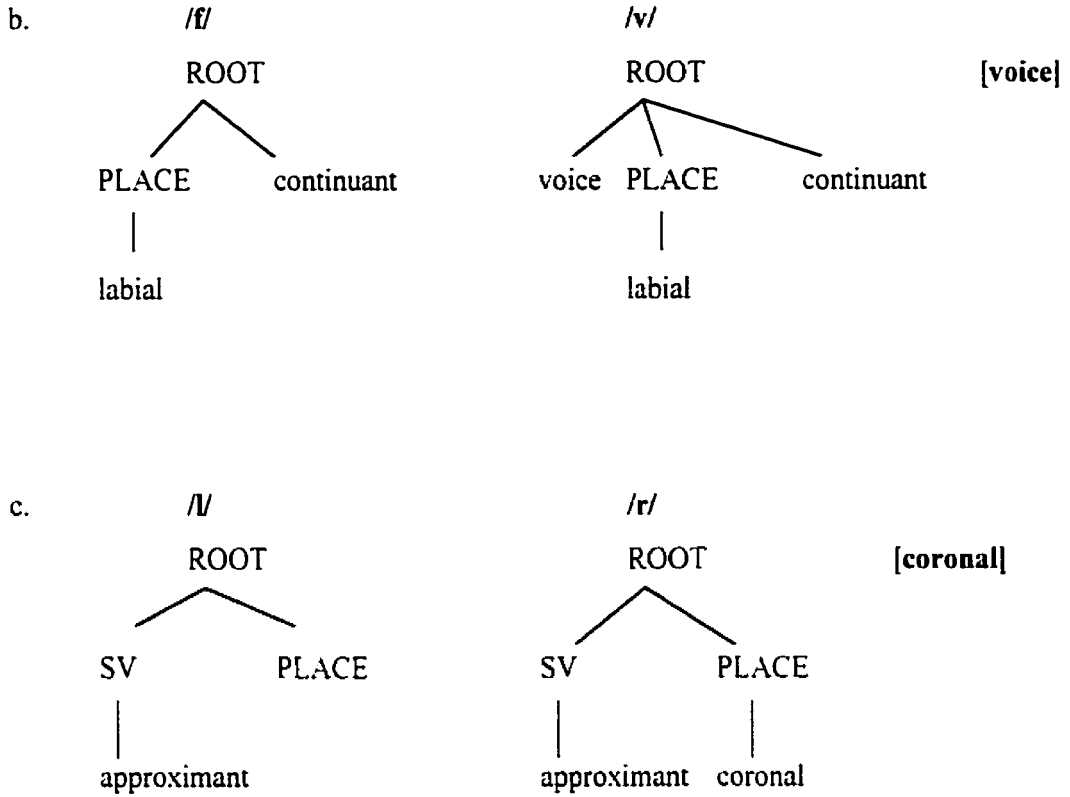
3.3.2 Experiment 2: Japanese speakers' acquisition of English /l, r/, /b, v/ and /f, v/

On the basis of the results discussed in Section 3.3.1 above, Brown ran a second experiment to test the hypothesis that perception of an L2 phonemic contrast is dependent on features present in the L1 rather than on the precise segmental inventory of the native language. Because features appear to provide the basis for acquisition of non-native phonemic contrasts, Brown tested Japanese speakers' acquisition of various English phonemic contrasts that differed on the basis of a single feature. Some features were required in the Japanese feature geometry for the representation of other segmental contrasts while others were not. (See (6) above for the fully specified, adult feature hierarchy for Japanese.)

The English contrastive pairs /l ~ r/, /b ~ v/ and /f ~ v/ were chosen as experimental stimuli because none of the three pairs contrast in Japanese and each pair differs in the level of difficulty it creates for the learner. Crucially, each of the contrasting pairs minimally differs on the basis of a single feature. In (7) I present the structural representations for each of the pairs with the contrasting feature indicated in bold to the right³:



³ The superordinate nodes SUPRALARYNGEAL and LARYNGEAL have been omitted for ease of exposition as they are not relevant to this discussion.



(Brown, 2000: 22)

Two of the three contrasts represented in (7) require a feature manipulated in Japanese for the representation of another phoneme while the third contrast does not: the English segments /b/ versus /v/ are distinguished by the feature [continuant] which Japanese requires to distinguish stop versus continuants such as /s/ versus /t/: the segments /f ~ v/ are distinguished by the feature [voice] which is required in Japanese to distinguish between the segments /p ~ b/: and the segments /l ~ r/ require the feature [coronal], which, as we saw in Experiment 1, Section 3.3.1, is not required in the representation of any segment in Japanese. Brown's claim is that Japanese speakers should be able to acquire those contrasts, specifically /b ~ v/ and /f ~ v/, that require a feature manipulated elsewhere in their native grammar.

To test this hypothesis, Brown tested thirty adult ESL students using the same two task described in Experiment 1 in Section 3.3.1. The stimuli included the /p ~ b/ contrast found in Japanese as a foil; also, the stimuli were limited to syllable onset position due to the increased number of contrasts to be tested. Results are presented in Table 3.3:

Table 3.3 Mean Results of Discrimination Task: English Phonemic Contrasts

	<i>Phonemic Contrast: Percent Correct</i>			
<i>Language</i>	<i>/l ~ r/</i>	<i>/b ~ v/</i>	<i>/f ~ v/</i>	<i>/p ~ b/</i>
Japanese	35.6%	90%	96.1%	95%

Results indicated that Japanese speakers' performance on the English /l ~ r/ contrast was significantly worse than performance on the other three contrasts, as predicted by the hypothesis. Performance on the other three contrasts /p ~ b/, /b ~ v/, and /f ~ v/, however, was near perfect. Moreover, the performance on these three contrasts was not significantly different from one another, as we would predict given that each of these three contrasts requires a feature manipulated elsewhere in the Japanese hierarchy.

As in Experiment 1, performance on the Forced Choice Picture Selection task supported the finding of the AX discrimination task. The subjects performance on the /l ~ r/ contrast was significantly poorer than performance on the other three contrasts, as predicted. Results are presented in Table 3.4:

Table 3.4 Mean Results of Picture Selection Task: English Phonemic Contrasts

	<i>Phonemic Contrast: Percent Correct</i>			
<i>Language</i>	<i>/l ~ r/</i>	<i>/b ~ v/</i>	<i>/f ~ v/</i>	<i>/p ~ b/</i>
Japanese	52.7%	88.9%	96.7%	93.9%

Results from this experiment thus appeared to support the proposal that the perception of non-native phonological contrasts is constrained by features, not segments.

present in the learners' native phonemic inventory. Even though learners lack acoustic, phonetic or phonemic experience with a non-native contrast, if they utilise the particular distinguishing feature to contrast other segments in their native inventory, then this experience should allow them to perceive other contrasts along the same dimension represented by that feature. The feature is in place to filter the incoming acoustic signal into the appropriate phonemic category.

We now turn to a further piece of evidence in support of Brown's feature theory of L1 interference in L2 perception: Matthews' (1997) experimental research showing the varying influence of pronunciation training on novel segmental contrasts that both are, and are not, characterized by features present in the L1.

3.4 *The Influence of Pronunciation Training on the Perception of Second Language Contrasts (Matthews, 1997)*

To address the broad question of whether second language learners can acquire novel segmental contrasts, Matthews investigated the effects of pronunciation training on the acquisition of segmental contrasts. Following Brown (1993, 1997), Matthews argues that the crucial factor in determining successful acquisition of novel segmental contrasts is the native language feature geometry which imposes a template within which the novel segmental contrasts are perceived. Since the critical factor in acquiring novel segmental representations is the features present in the native language feature hierarchy, Matthews argues that the learner will only be sensitive to those non-native contrasts that are distinguished along dimensions corresponding to features in the native feature geometry: if a novel phonemic contrast requires a contrastive feature that is manipulated elsewhere in the geometry, then the learner will be able to perceive its contrastive use in the input and will subsequently be able to acquire the novel representation. Conversely, if a feature is not represented anywhere in the native structure, perception -- and hence production -- of that contrast will be precluded.

To test this hypothesis, Matthews studied the effects of pronunciation training on Japanese speakers' production of various English segmental contrasts. He noted that Japanese and English have segmental contrasts other than the much studied /l/ vs /r/

contrast, and divided the contrasts into three categories: those contrasts in which both members are present in Japanese; those contrasts where one member of the pair is present in Japanese while the other is absent; and those contrasts where both members are absent from Japanese. The contrasts are presented in Table 3.5:

Table 3.5 Relationship Between English Contrasts and Japanese Inventory

<i>English Segmental Contrast</i>	<i>Present in Japanese</i>	<i>Absent from Japanese</i>
[p] ~ [b] [t] ~ [s]	[p] [b] [t] [s]	
[b] ~ [v] [s] ~ [θ]	[b] [ɰ]	[v] [θ]
[θ] ~ [f] [l] ~ [r]		[θ] [f] [l] [r]

Matthews argued that the presence or absence of the segment in the native inventory was not enough to determine the positive influence of pronunciation training on perceptual categories; rather, it was the feature representation established in the course of L1 acquisition. The results of the experiment supported this argument: contrasts indicated in bold shadow outline in Table 3.5 showed improvement after pronunciation training. Post-hoc Scheffe F-tests indicated significant differences between pre- and post-test scores for the [b] ~ [v] contrast and the [θ] ~ [f] contrasts ($F(65,11) = 1.037$; $p = .0001$). While the pretest to post-test improvement on the [s] ~ [θ] contrasts resembles that of both the [b] ~ [v] and the [θ] ~ [f] contrasts, the difference was not significant. Both the [p] ~ [b] contrast, which does occur in Japanese, and the [l] ~ [r] contrast, which does not, showed only negligible change.

Matthews concluded that the L1 phonological system constrains L2 development of novel segmental categories and that instruction in the pronunciation of non-native

segmental contrasts is effective in establishing novel segmental representations which can then be used to perceptually contrast novel segments if and only if the non-native contrast is characterised by a feature present in the L1 inventory.

In this chapter I have outlined the requirements for a good theory of L2 acquisition. I then presented one such model that fulfills all three requirements: Brown's model of L1 interference. This is the model adopted in this thesis. We have also looked at research supporting the argument that it is the featural rather than the segmental level determining L2 perception, and subsequent production of novel L2 sounds. Having looked at the field of L2 acquisition in a broad sense, it is time to move to the specifics of English speakers' acquisition of the Czech sound system. In the following chapter I present the phonological hierarchies of Czech and English, followed by Chapter Five in which I present the experimental data and analysis.

CHAPTER FOUR

SEGMENTAL PHONOLOGY OF CZECH AND ENGLISH

0.0 BACKGROUND

This chapter establishes the Feature Geometries for Czech and English as a framework for the experimental data on English speakers' acquisition of the Czech alveolar versus palatal stop contrast examined in Chapter Five. Under the phonological model of L1 interference assumed in this thesis, it is the L1 feature inventory that dictates whether or not speakers of a particular language will be able to acquire a new phonemic contrast in the L2. Thus, in order to determine whether or not English speaking learners of Czech will have difficulty in acquiring the novel Czech segmental stop contrast of alveolar /t, d/ versus palatal /c, j/, we must first establish the language-specific featural inventories of these two languages.

Chapter Four begins with an overview of the segmental phonology of Czech with a focus on alveolar and palatal plosives. We then turn to the articulatory and acoustic characteristics of alveolar and palatal segments, followed by a discussion of the phonological features required for to represent the contrast between alveolar and palatal segments and the phonemic representations for Czech /t, d/ and /c, j/. Finally, in Section 3 we look at the segmental inventory of English, focusing on the feature specification of the three Coronal places of articulation for English fricatives.

1.0 THE SEGMENTAL INVENTORY OF CZECH

Czech is a member of the Western branch of the Slavic family, spoken by about 10 million native speakers in the Czech Republic (Short, 1993). There are also fairly large Czech-speaking communities in North America and small, isolated pockets in neighboring European countries. The two closest related languages are Polish and Slovak, which is mutually intelligible with Czech.

1.1 Vowels

The Czech vowel system consists of five short vowels /ɪ, ɛ, a, o, u/ and a corresponding set of five long vowel counterparts /i:, ɛ:, a:, o:, u:/. Short and long vowels contrast in all positions. There are also three falling diphthongs /ou/, /au/ and /ɛu/. With the exception of the pair /ɪ/ and /i:/, the quality of short and long vowels differs only slightly. In the case of /ɪ/ and /i:/, the short vowel is more central and substantially less close than the long vowel /i:/ (Dankovičová, 1999).

1.2 Consonants

As this thesis is primarily concerned with the acquisition of the segmental contrast between the alveolar and palatal stop consonants of Czech by native English speakers, I will not go into details of each segment in the Czech inventory. Instead, I present a chart of the segmental inventory of Czech followed by a brief overview of Czech plosives, before getting to the heart of the matter, alveolar versus palatal stops. Table 4.1 below presents the segmental inventory of Czech consonants. For each place of articulation, the voiceless member of a pair is presented on the left and the voiced member to the right:

Table 4.1 Phonemic Inventory of Czech Consonants (Dankovičová, 1999)

	Bilabial	Labiodental	Alveolar	Postalveolar	Palatal	Velar	Glottal
Plosive	p b		t d		c ɟ	k ɡ	
Nasal	m		n		ɲ		
Fricative		f v	s z	ʃ ʒ		x	h
Affricate			ts	tʃ			
Trill			r				
Trill Fricative			ʀ				
Approximant					j		
Lateral Approximant			l				

1.2.1 Plosives

All Czech stops occur in voiced-voiceless pairs and are not normally aspirated. Czech contrasts four places of articulation for each voiceless-voiced pair of stops: bilabial /p, b/, alveolar /t, d/, velar /k, ɡ/, and palatal /c, ɟ/. As will be discussed in detail, a key difference between the Czech and English consonant inventories is in the palatal area of articulation. Czech has a phonemic contrast between two Coronal places of articulation: alveolar /t, d, n/, which do occur in English, versus palatal /c, ɟ, ɲ/, which do not. To acquire this new segmental contrast, then, English speakers must (i), be able to articulate palatal sounds and (ii), establish new segmental representations for the novel phonemic contrast between alveolar and palatal.

Although we are concerned here with phonology and not orthography, a note is needed on the orthographic conventions of Czech before I present tokens of minimal pairs contrasting alveolar /t, d/ or palatal /c, ɟ/. Czech has a highly consistent sound - spelling correspondence where the orthography of words containing either alveolar /t, d/ or palatal /c, ɟ/ is determined by the quality of the following vowel. To accustom the reader to

these conventions before I present the experimental data, as well as to show that alveolars and palatals can precede any vowel, I will present five sets of minimal pairs contrasting alveolars and palatals. Minimal pairs will be presented with the following vowels in alphabetical order: 'a', 'e', 'i', 'o', 'u'.

In Czech, both alveolar /t, d/ and palatal /c, j/ are represented orthographically by the letters 't' and 'd'; however, when the stop is a palatal then it is indicated by means of a diacritic or a different vowel symbol, depending on the quality of the following vowel. When the following vowel is an 'a', the palatal sound is indicated by an apostrophe after the 't' or 'd'. In (1) I present minimal pairs contrasting alveolar /t, d/ and palatal /c, j/ with a following 'a'; voiceless /t~ c/ is given in (1a-b); voiced /d ~ j/ is in (1c-d):

(1)	Orthographic Form	Phonetic form	gloss
a.	vráta	[vr̩a:ta]	'gate'
b.	Vrát'a	[vr̩a:ca]	'(proper name)'
c.	dás	[da:s]	'gum'
d.	d'as	[jas]	'deuce'

When the following vowel is an 'e', the palatal sound is indicated by means of the Czech diacritic 'háček' over the 'e': that is, 'ě'. Minimal pairs contrasting alveolar /t, d/ and palatal /c, j/ with a following 'e' are given in (2); voiceless /t~ c/ is in (2a-b); voiced /d ~ j/ is in (2c-d):

(2)	Orthographic Form	Phonetic form	gloss
a.	teka	[teka]	'run' (3.p.s)
b.	těka	[ce:ka]	'wander' (3.p.s)
c.	dekovat	[dekovat]	'to steal'
d.	děkovat	[jekovat]	'to thank'

With respect to the following vowel ‘i’, the pattern is slightly different than for either ‘a’ or ‘e’ in that a different character rather than a diacritic is used to indicate the palatal stop. When the following vowel is an ‘i’, the palatal sound is indicated by the character ‘i’ while alveolar ‘t’ or ‘d’ are followed by the letter ‘y’. It is important to note that the vowel quality does not change: the use of a different vowel character indicates only that the preceding consonant is either alveolar or palatal. In (3) I present minimal pairs contrasting alveolar /t, /d and palatal /c, j/ with a following ‘i / y’: voiceless /t~ c/ is given in (3a-b); voiced /d ~ j/ is in (3c-d):

(3)	Orthographic Form	Phonetic form	gloss
a.	tyka	[tka]	‘(to) tick’
b.	tika	[cika]	‘(to) spy’
c.	mlady	[mladi:]	‘young (adj.)’
d.	mladi	[mlaji:]	‘youth’

When the following vowel is either an ‘o’ or a ‘u’, the palatal sound is indicated by an apostrophe after the ‘t’ or ‘d’ as we saw for the vowel ‘a’. In (4) I present minimal pairs contrasting alveolar /t, /d and palatal /c, j/ with a following ‘o’: voiceless /t~ c/ is given in (4a-b); voiced /d ~ j/ is in (4c-d). Minimal pairs contrasting alveolar /t, /d and palatal /c, j/ with a following ‘u’ are given in (5); voiceless /t~ c/ is in (5a-b); voiced /d ~ j/ is in (5c-d):

(4)	Orthographic Form	Phonetic form	gloss
a.	topý	[topi:]	‘heating’
b.	t’opý	[copi:]	‘trotting’
c.	dobu	[dobu]	‘time’
d.	d’obu	[jobu]	‘picking’ (l.p.s)’

(5)	Orthographic Form	Phonetic form	gloss
a.	tuk	[tuk]	'fat' (noun)
b.	t'uk	[cuk]	'tap' (noun)
c.	dub	[dub]	'oak'
d.	d'ub	[jub]	'nudge' (noun)

Word-finally, the palatal stop is also indicated by means of an apostrophe.

Examples of word-final alveolar /t ~ c/ are given in (6a-b). No tokens of /d ~ j/ are given as voiced stops devoice word-finally in Czech.

(6)	Orthographic Form	Phonetic form	gloss
a.	plet	[plet]	'to knit'
b.	plet'	[plec]	'complexion'

At this point, it should be clear that Czech contrasts alveolar versus palatal stops. In subsequent sections, I will argue that the Coronal place node in Czech requires more complex elaboration of the Place node to represent these two contrasting coronal places of articulation. A bare unspecified Place node for the unmarked coronal segment would be insufficient to distinguish between the two segments as the contrastive distinction would be conflated. However, before we establish which feature is required to elaborate the Coronal Place node to represent this contrast, let's first look at the articulatory and phonetic characteristics of alveolar and palatal sounds in Czech before taking a closer look at the features used to represent them.

2.0 CZECH CORONALS

As we saw in Chapter One, coronal segments are sounds which are produced with the front part of the tongue including the tongue tip and blade (Maddieson, 1984). Czech

has two places of Coronal articulation, alveolar and palatal, which we will look at in some detail in this section.

2.1 *Alveolars*

Alveolar segments are those sounds which have their point of constriction on the alveolar ridge just behind the upper front teeth. Alveolar sounds are classified as [anterior] sounds along with linguolabials, interdentals, and dentals (Keating, 1991); the class of anterior sounds may vary in terms of the apicality. Keating (1991:33) points out that in both French and English, there is inter-speaker variation in both place and manner for dental and alveolar sounds.

2.2 *Palatals*

2.2.1 *Palatal articulations*

Palatal sounds are made with the point of constriction behind the alveolar ridge on the hard palate. The hard palate, however, is quite a large area with a corresponding large number of possible constriction points. As palatal sounds can vary as to relative frontness or backness with respect to constriction, no consensus exists among phoneticians as to the precise articulatory specification.

Due to the variety of articulatory constrictions subsumed under the name 'palatal', and based on X-ray and palatographic data showing that palatal consonants can be more finely controlled than previously thought, Recasens (1990) expanded the traditional two palatal zones, palato-alveolar and palatal, into four classes: alveolo-palatals, front palatals, mid palatals and back palatals. The term palatal serves as a cover term for all four of these classes. Looking at palatograms (Chlumsky, 1941; Hála, 1923, 1962) for the voiceless Czech palatal stop /c/, Recasens notes that complete contact occurs at the postalveolar and prepalatal zones, which is substantially more front than for other palatal sounds. As a result, he calls Czech /c, ɟ/ alveolo-palatal as does Hall (1996).

Keating & Lahiri (1993) looked at palatograms & linguograms to compare front velars, palatalised velars and palatals. They found that Czech palatals had a wide occlusion, and that the central contact of this occlusion was made with the tongue blade.

and not the tongue tip or body: Czech coronals looked like a long coronal stop.

In sum: the work by both Recasens (1990) and Keating & Lahiri (1993) shows that, phonetically, Czech palatal segments are coronal places of articulation. In the next section we turn to phonological evidence to further support this claim.

2.2.2 *Palatals as Coronal sounds*

We now turn to the phonological characteristics of palatal segments and the features used in their representation. As background, I will start with a brief overview of features traditionally used to characterise palatal segments before stating my position in Section 2.3.

The original feature theory (see Jakobson, 1938/1962; Jakobson, Fant & Halle, 1952) used the binary acoustic feature [grave] to differentiate between [+grave] peripheral sounds such as /p, b/ and /k, g/ versus [-grave] non-peripheral sounds such as /t, d/.

This binary acoustic feature [grave] was replaced by Chomsky and Halle in the SPE by the binary articulatory feature [coronal]. Chomsky & Halle claimed that palatal segments are non-coronal, or [-coronal] sounds, since they were considered to involve tongue-body articulations. Under this definition, dental, alveolar, retroflex, palato-alveolar, and alveolo-palatal segments are grouped together as [+coronal] while labial, palatal and velar sounds are classified as [-coronal].

However, the classification of palatal sounds as [-coronal] proved controversial as it fails to capture the fact that alveolars such as /t, d, n/ and palatals such as /c, ʃ, ɲ/ often pattern together phonologically. For example (as discussed in Chapter One, Section 3.2), in Sanskrit an alveolar [n] is retroflexed to [ɳ] if it follows a retroflex continuant as long as there is no intervening alveolar, retroflex or palatal sound (Whitney, 1885; Odden, 1978). Similarly, in Hungarian the coronal fricatives [s, z, ʃ, ʒ] assimilate to [ts, dz, tʃ, dʒ] respectively when they are preceded by either alveolar [t, d] or palatal [c, ɟ] (Vago, 1989; Hume, 1992).

To better reflect the common patterning of alveolars and palatals, palatals were redefined by post-SPE phonologists as coronal sounds. Initially, post-SPE phonologists argued *against* the replacement of [grave] in favor of [coronal] since the feature [-grave]

can capture the natural class of alveolars and palatals (see Vennemann & Ladefoged, 1973; Vago, 1976; and Odden, 1978). A resolution was reached by redefining palatals as Coronal sounds on the basis of the articulatory affinity between alveolar /t, d, n/ and palatal /c, ʃ, ɲ/ (see for e.g. Lahiri & Blumstein, 1984; Keating 1988, 1991). This redefinition of defining features was prompted by Pagliuca & Mowrey (1980) who noted that the phonological patterning of alveolars with palatals has an *articulatory* motivation (thus eliminating the need for the *acoustic* feature grave).

Currently, most phonologists are in agreement that palatal sounds are indeed coronal. The classification of palatals as coronal is justified on several bases. First, palatals pattern phonologically with alveolars as we saw above in the Sanskrit and Hungarian examples. Since partial motivation for feature based analyses is to represent natural classes, alveolars and palatals require a common feature to capture this pattern. Secondly, palatal sounds in general are articulated very far forward in the mouth. Keating (1991: 38) notes that “palatals are articulated much further forward in the mouth, and on the tongue, than has often been assumed” so that theoretically palatals are ‘next to’ velars, but rather far apart in practice. Thirdly, and crucially, in the articulation of palatals the tongue blade touches just behind alveolar ridge so that the point of constriction itself is coronal.

2.3 *Czech Feature Hierarchy*

We have seen that Czech contrasts two stop places of articulation that are phonologically coronal: alveolar /t, d/ versus /c, ɟ/. Given, then, that Czech has two coronal places of articulation, the Place node Coronal needs to be elaborated for one of the segments in order to distinguish between the two segments. We now need to establish and motivate the dependent feature of the Coronal node. The question is: What is the contrasting feature between alveolar and palatal places of articulation? Stevens (1972, 1989) points out that some articulatory gestures are easier to make than others for physiological reasons, and that considerations of ease of articulation and auditory distinctiveness can influence the phonetic structure of a language. He notes that this may account for the comparative lack of palatal sounds among the world’s languages. If we

assume that marked sounds are less frequent, then palatals are more marked than alveolars.

Moreover, the historical development of palatal segments in Czech would support the assumption that palatals are more marked than alveolars. At some point, Proto-Slavic had three places of articulation for stops: /p, t, k/. However, the palatalization of alveolar consonants occurred when an alveolar was followed by a front vowel (Carlton, 1991).

Rice & Avery (1991) have shown that marked sounds are structurally more complex. In the phonological framework of feature geometry, segment complexity is seen in as a more elaborate feature structure. Thus, we would expect palatal segments to have a feature structure with more featural representation than alveolar segments.

Therefore, based on markedness and historical considerations I will argue that palatal, rather than alveolar, segments in Czech are more structurally complex and require further elaboration of the Place node by a dependent feature. In the next section we discuss the feature [anterior], commonly used to distinguish between anterior and posterior coronal segments, as well as its converse, [posterior]. I then present arguments for the representation of the marked palatal segments by the dependent feature [posterior].

2.3.1 *Anterior versus posterior*

Assuming that palatals are more marked than alveolar segments and consequently require further elaboration of the coronal node, we now need to determine which feature is contrastive in the representation of these two coronal segments. I will present an overview of features that have traditionally been used to distinguish coronal segments before I present arguments for the feature I assume in this thesis.

In the SPE, Chomsky & Halle used the binary features [+anterior] and [-anterior] (replacing the feature [compact]) to distinguish between anterior and posterior coronal segments in languages that contrast more than one coronal place of articulation. The feature [+anterior] described sounds produced at or in front of the alveolar ridge such as labial, interdental, dental and alveolar sounds; these are distinguished from [-anterior] segments produced behind the alveolar ridge such as palato-alveolar, retroflex, alveolo-palatal, palatal, velar, and uvular.

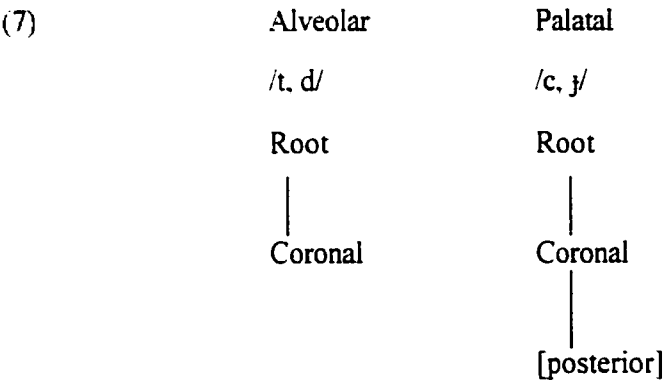
Since Sagey (1986), however, the use of [anterior] in Feature Geometry theory has for the most part been limited to being a terminal node for the coronal place of articulation only, and is thus used primarily to distinguish more front coronal segments such as /s, z/ from more back coronal segments such as /ʃ, ʒ/. One possibility, then, would be to define Czech alveolar segments as [anterior] and palatal segments as [-anterior]. However, as discussed in Chapter 1 on phonological theory and feature geometry, recent empirical and theoretical evidence has argued that all features are monovalent (see Anderson & Ewen, 1987; Avery & Rice, 1989; van der Hulst, 1989), which is the stance I take in this thesis. I argue that articulator features are monovalent, or privative, with contrasts represented with the presence or absence of a node rather than by binary [+] or [-] values.

Assuming a monovalent feature [anterior] would force us to define the more front alveolar segments as [anterior], while palatal segments would be less elaborated structurally and would be captured by the feature coronal, without further elaboration by a dependent feature. However, representing the marked palatal segment with less structural elaboration than alveolar segments is undesirable, given that marked sounds are considered to be structurally more complex.

To avoid the undesirable consequence of having a more marked segment less structurally elaborate than the unmarked segment, I propose that the converse of the feature [anterior] be used. That is, the place feature [posterior] will be used as a dependent of the Coronal node to characterise the marked palatal sounds so that they are structurally more elaborate than alveolar segments. This has the desired result of marking structural complexity directly in the representation. Unmarked sounds, in this case alveolar segments, need not be overtly represented in the featural structure but are instead defined by the absence of the opposite marked feature, [posterior]. Thus, I propose that the Czech palatal /c, ɟ/ require elaboration of the coronal node by the feature [posterior], while the alveolar pair /t, d/ are characterised by only the articulator feature Coronal².

² Alternatively, it has been suggested that the contrastive feature could be [distributed]. The feature [distributed] refers to manners of articulation with reference to the length of constriction in the airflow and differentiates between apical sounds (made with the tongue tip) and laminal sounds (made with the tongue blade). I avoid the use of this feature as, however, as it is a constriction based feature whereas [posterior] is an articulator feature.

Alveolar segments would be interpreted by the absence of the dependent feature [posterior]. The representations for Czech palatal and alveolar segments are given in (7):



The representations in (7) correspond to markedness relations found cross-linguistically between alveolars and palatals. Now that we have the defining feature representing the contrast between alveolar and palatal stops in Czech, we can make a hypothesis as to English speakers’ acquisition of the Czech alveolar versus palatal contrast:

Hypothesis 1: If English does represent the feature [posterior] somewhere in its inventory, then English speakers should be able to perceive the phonemic stop contrast between Czech alveolar /t, d/ and palatal /c, ʃ/, and hence, will be able to establish new segmental representations.

As we will see in the next section, I argue that English does in fact represent the feature [posterior] in its inventory.

3.0 THE PHONOLOGY OF ENGLISH

We now turn to the segmental inventory and phonological feature structure of English. As with the Czech inventory, I will not go into details of each segment in the English inventory; instead, I present a chart of the segmental inventory followed by an

overview of English fricatives. Table 4.2 below presents the segmental inventory of English consonants. For each place of articulation, the voiceless member of a pair is presented on the left and the voiced member on the right:

Table 4.2 Phonemic Inventory of English Consonants

	Bilabial	Labio-dental	Inter-dental	Alveolar	Alveo-palatal	Palatal	Velar	Glottal
Plosives	p b			t d			k g	
Fricatives		f v	θ ð	s z	ʃ ʒ			h
Affricates					tʃ dʒ			
Nasals	m			n			ŋ	
Approximant	w			r		j		
Lateral Approximant				l				

As discussed extensively in Section 1, this chapter, English has a three-way stop contrast in comparison to the four-way stop contrast in Czech. Because English has a single Coronal place of articulation for stops, it requires only a bare Place node to represent this contrast. Further elaboration is redundant: English does not require the dependent feature [posterior] to characterise stops as does Czech. However, as we will see in the next section, English does contrast three places of Coronal articulation for fricatives: the alveolar /s, z/, the alveo-palatal /ʃ, ʒ/, and the dental /θ, ð/.

3.1 English Coronal Fricatives

English contrasts three places of articulation for each voiceless-voiced pair of fricatives: alveolar /s, z/, alveo-palatal /ʃ, ʒ/, and dental /θ, ð/. Phonologically speaking, all three places of articulation are Coronal places of articulation so that English contrasts three coronal fricatives. Clearly, Coronal as place of articulation is not specific enough to

distinguish between them: some other feature(s) must be used. Thus, further elaboration of the Coronal node is essential as a bare Place node of Coronal would conflate the phonemic distinction. Which segmental pair is more structurally elaborate? In the next section, I present evidence for my claim that alveolar /s, z/ are the default, underspecified phonemes, and that alveo-palatal /ʃ, ʒ/ and dental /θ, ð/ require further elaboration of the Place node by a dependent feature.

3.2 *Feature Representations of English Fricatives*

A review of a number of longitudinal and detailed cross-sectional studies of English speaking children's order of acquisition of phonemes (see for e.g. Templin, 1957; Prather, Hedrick & Kern, 1975; Arlt & Goodban, 1976) reveals three consistent patterns with respect to the English coronal fricatives alveolar /s, z/, alveo-palatal /ʃ, ʒ/ and dental /θ, ð/: (i), there is a consistent order to children's acquisition of English coronal fricatives, with the alveolar /s/ being acquired first followed by the alveo-palatal /ʃ/ and, somewhat later, the (inter)dental /θ/: (ii), the fricatives /ʃ/ and /θ/ are among the last sounds to be acquired, and (iii), /ʃ/ and /θ/ are among the most common sounds mis-produced by children with language delays (Grunwell, 1982; Stoel-Gammon & Dunn, 1985; Vihman et al., 1986)³.

Grunwell (1982) synthesized data from a number of first language acquisition studies in order to delineate a set of stages indicating which phonemes are expected at each stage. She points out that the alveolar /s/ is the first of the coronal fricatives to be acquired (fairly early) at approximately age 3;0. At this stage, the contrast between /s/, /ʃ/ and /θ/ is neutralised, so that a child may produce [sip] instead of the adult form [ʃip] for 'sheep' and [bæs] instead of [bæθ] for 'bath'.

The second coronal fricative to appear is the alveo-palatal /ʃ/, which is not

³ It should be noted that these studies represent developmental norms based on large numbers of children; individual children may vary in their production. However, Jakobson (1971) notes that while the particular age of acquisition of a particular phoneme may vary, the general order of acquisition is quite consistent across children.

acquired until approximately ages 3;6 to 4;6.

Finally, the last coronal fricatives to be acquired are the (inter)-dental /θ/ and /ð/, which appear quite late at around age 4;6 but are not often not acquired consistently until age 5 or 6.

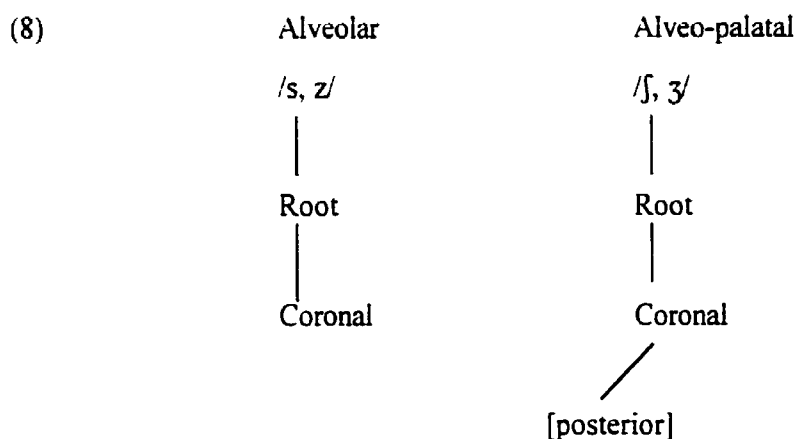
Under the segmental acquisition as structural elaboration framework argued for in this thesis, the three patterns above can be attributed to the underspecification of /s, z/ for Place. As discussed extensively in Chapter Two, the acquisition of segmental structure can be seen as a process of expanding the initial, basic feature structure provided by UG to the language-specific feature geometry used to differentiate native phonemes by adult speakers. Thus, structurally less complex (default) phonemes are acquired before those requiring further featural elaboration. The alveolar segments /s, z/ are acquired before either alveo-palatal /ʃ, ʒ/, or dental /θ, ð/ as alveolar /s, z/ require less structural elaboration. I argue that alveolar /s, z/ is the underspecified default phoneme, specified only for Coronal (the coronal node itself being triggered by contrast), while both alveo-palatal /ʃ, ʒ/, and dental /θ, ð/ require further elaboration of the Place node to differentiate them from /s, z/, and from each other. Because the initial structure is minimally specified and elaboration (prompted by detection of phonemic contrast in the input) is necessary for contrast, the segments /ʃ/ and /θ/ are acquired somewhat later than /s/.

Which features are required to distinguish alveo-palatal /ʃ, ʒ/ and dental /θ, ð/ from /s, z/, and from each other? A number of features have been argued for as being the key feature distinguishing between these English coronal fricatives, including [strident], [grave], [anterior] and [distributed]. Moreover, the debate does not stop at the level of deciding precisely which feature is the contrastive feature: phonologists differ even in ascribing a consistent value to each segment. I will not go into details of each of these proposals: instead, I present arguments for the feature [posterior] as representing alveo-palatal /ʃ, ʒ/ and [distributed] as the contrastive feature separating dental /θ, ð/ from the other two coronal fricatives.

3.2.1 *Feature Representation of Alveo-palatal /ʃ, ʒ/*

As discussed in Section 2.3.1, Chomsky & Halle (1968) used the binary features [+anterior] and [-anterior] to distinguish between anterior and posterior coronal segments, pointing out that anterior sounds have a constriction before the palato-alveolar region of the mouth while nonanterior sounds are produced without such a constriction. This proposal was rejected in this thesis as features are assumed to be privative. Thus, we cannot use these binary features to distinguish between the more front English /s/ and the more back /ʃ/.

In the same section, I argued for structural representations of markedness by representing the more marked Czech palatal with the dependent of the Coronal node, the feature [posterior]. On the same grounds, I propose that because English alveo-palatals /ʃ, ʒ/ are acquired later (and are thus more marked segments), they are more structurally complex than alveolar /s, z/ and must be represented by the elaborated Coronal Place node. That is, representing the English alveolar fricatives /s, z/ with a dependent feature [anterior] would result in an unmarked segment with more structural elaboration than the marked segment. I argue that it is the marked alveo-palatal /ʃ, ʒ/ which requires further structural elaboration. Alveolar /s, z/ as the default segments need not be overtly represented in the featural structure but are instead defined by the absence of the opposite marked feature, [posterior]. Thus, I propose that the English alveo-palatals /ʃ, ʒ/ require elaboration of the coronal node by the feature [posterior], while the alveolar pair /s, z/ are characterised by only the articulator feature Coronal. Representing English alveo-palatals /ʃ, ʒ/ with the feature [posterior] marks structural complexity directly in the representation. The representations for English alveolar and alveo-palatal segments are given in (8):



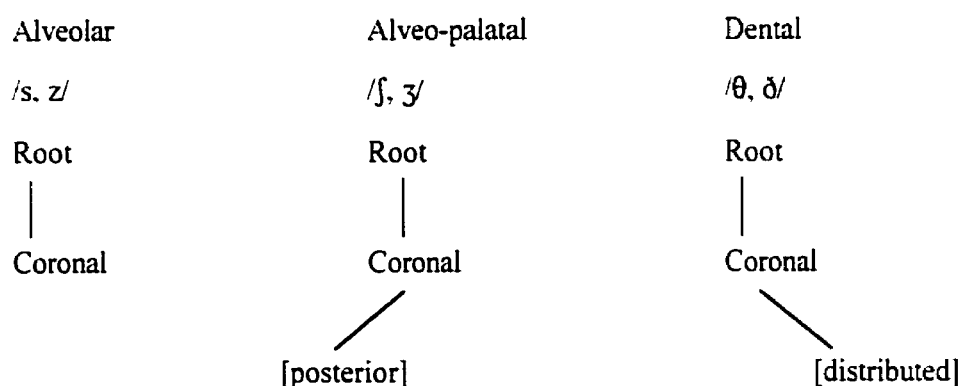
3.2.2 Feature Representation of Dental /θ, ð/

The SPE defines the feature [distributed] as a manner property differentiating apical (pronounced with the tongue tip) and laminal (articulated with the tongue blade) manners of articulation with reference to the length of constriction in the airflow: in fact, the term apical-laminal is sometimes used in place of [distributed] (Clements, 1989). Sounds that are [-distributed] have shorter constriction and (usually) apical articulations whereas [+distributed] sounds have longer constrictions and laminal articulations. Many post-SPE researchers have limited [distributed] to the coronal place of articulation only, making the feature less controversial. However, the debate is ongoing as to which pair of coronal fricatives, alveolar /s, z/ or dental /θ, ð/, are [+distributed]; the debate hinges on the varying definitions of the term distributed with those claiming that /s, z/ are [distributed] focused on its apical production and not the length of the constriction, while those arguing for /θ, ð/ as [distributed] focus on the extended length of the constriction as well as the apical production.

Under the Underspecification theory of Selkirk (1988, 1993), the feature [distributed] refers to degree of stricture, which is the position I take in this thesis in arguing for /θ, ð/ as [distributed] due to the extended constriction of the active articulator. I argue that (inter)dental /θ, ð/ contrast with /s, z/ on the basis of the feature [distributed] as a dependent of the Coronal node. The feature representations for the three English

Coronal fricatives are in (9):

(9)



We have seen that English contrasts only one place of articulation for stops, thereby making further elaboration of the featural structure redundant. However, English does contrast three coronal places of articulation for fricatives and thus requires further elaboration of the structural representation for one of these segmental pairs. Following work in child language acquisition I argue that the later acquired segments alveo-palatal /ʃ, ʒ/ and dental /θ, ð/ are more structurally complex and required further elaboration of the Place node. Following arguments for Czech, I argue that the marked alveo-palatal /ʃ, ʒ/ was represented by the dependent feature [posterior], while the marked dental segments /θ, ð/ required the feature [distributed]. English thus does represent the feature [posterior] somewhere in its inventory. Under Brown's model of L1 interference in L2 segmental acquisition, if an L2 learner manipulates the feature representing a particular phonemic contrast somewhere in the inventory, they should be able to perceive any contrast requiring that feature. In other words, the contrasting feature need not be utilised for the same contrast in the L1 as in the L2. Thus, although Czech requires the feature [posterior] to maintain a distinction between alveolar and palatal stops which is unnecessary in English, English speakers do represent the feature [posterior] to distinguish between alveolar and alveo-palatal fricatives. Consequently, English speakers should be able to perceive the distinction between alveolar and palatal, and hence, will be

able to establish new segmental representations.

In sum: this chapter establishes the feature geometries for Czech and English as a framework for the experimental data on English speakers' acquisition of the Czech alveolar versus palatal stop contrast examined in Chapter 5. I argued that English does in fact contrast the dependent feature [posterior], albeit for a different phonemic contrast than in Czech. However, even though English and Czech require the dependent feature [posterior] for different phonemic contrasts, perception of the novel L2 contrast should occur. Having established the feature inventories for Czech and English, we are now ready to turn to experimental data on English speakers' perception and production of the Czech alveolar versus palatal stop contrast.

CHAPTER FIVE

THE ACQUISITION OF CZECH PALATAL STOPS

0.0 Background

At this point in the thesis, we have covered the theoretical assumptions of generative phonology and both first and second language acquisition upon which I have based the experimental research on English speakers' acquisition of the Czech palatal versus alveolar stop contrast. Now that the background and motivation for the experimental design and research has been established, it is time to look at the experimental data itself. In this chapter, I outline the two experiments and results of native English speakers' perception and production of Czech palatal stops.

Two types of tasks were used to determine whether or not English speakers had acquired the non-native phonemic stop contrast of Czech alveolar /t, d/ versus palatal /c, ɟ/: production and perception tasks. Both types of tasks are required as reliance on either perception or production data to the exclusion of the other can be misleading in the area of adult second language acquisition. While production tasks can provide evidence that a learner has acquired the phonemic structure, there are pitfalls inherent in relying solely on production tasks for adult language learners, as adult learners, unlike children, have a fully developed motor control system and so can often produce the segment in question without having developed a mental (phonological) representation for the contrast. Experimental evidence has shown (e.g. Flege, 1995; Goto, 1971; Sheldon & Strange, 1982) that some adult second language learners can appropriately articulate the novel sounds in particular tasks while being unable to *perceptually* contrast the novel phonemes. For example, Goto (1971) found that Japanese subjects' production ability for the English /l/ versus /r/ contrast exceeded their perceptual abilities to distinguish the two phonemes, even for their own utterances: Sheldon and Strange (1982) replicated and extended these findings. By relying solely on production data, then, we may assume that the learner has acquired the appropriate phonological structure when in fact he or she has not.

Similarly, reliance solely on perceptual data can be misleading as some learners

are unable to produce a novel segment despite accurately perceiving the phonemic contrast. In order to avoid under- or over-assessing a learner's competence, then, it is necessary to look at perception data as well as at production data.

1.0 EXPERIMENT ONE: FORCED CHOICE PHONEME SELECTION TASK

The first experiment was the Forced Choice Phoneme Selection (henceforth FCPS) task. The FCPS task is designed to investigate learners' perceptions of non-native contrasts in which one member of a contrasting pair is found in the L1, while the other member is not. In this case, the FCPS task was used to test native English speakers' perceptual abilities in distinguishing the Czech segmental contrast between the alveolar stops /t, d/, which do occur in English, versus the palatal stops /c, ɟ/, which do not. Thus, the motivation for this experiment is to establish whether or not English speakers can *perceptually* contrast palatal and alveolar stops, or if they will tend to perceive the non-native palatal stops as tokens of alveolar stops).

There is abundant research in cross-linguistic perception showing that L2 learners make perceptual reference to existing L1 categories in attempting to impose structure on L2 speech (see for e.g. Bluhme (1969); Schouten (1975); Strange (1992)). Much of this research is in the area of voice onset time (VOT), which, in many languages (including English), is an acoustic cue used to distinguish between word-initial stop consonants. For example, Abramson & Lisker (1970) showed that L2 learners separate a VOT-varying continuum of stimuli into categories corresponding to the VOT of stop consonants of their L1. Similarly, Flege (1987, 1988, 1990) has shown that L2 learners project their L1 phonetic categories whenever possible on the sounds of the L2 in a process of "equivalence classification": only if the novel sound cannot be fit into an existing L1 category will a new phonetic category be constructed.

Given that production of a phonemic contrast depends in part upon accurate perception of a contrast (the learner must be aware that two sounds are in contrast), and that L2 learners appear to make perceptual reference to existing L1 categories to impose structure on novel L2 sounds, it was important to first establish English speakers' level of perceptual acuity for the Czech phonemic stop contrast of alveolar /t, d/ versus palatal

/c, ɟ/ before looking at their productive competence. Under the phonological model of L1 interference developed throughout this thesis, learners of a second language should be able to perceive a novel segmental contrast if that contrast is distinguished by a feature utilised in the native inventory for some phonemic contrast. If the native language does not utilise a particular feature on which the novel segmental contrast is based, then perception of that contrast should be precluded. The absence of the distinguishing feature in the L1 inventory means that both the novel segment as well as the familiar segment will be funneled into a single L1 phonemic category. In the case of English speakers learning Czech, the palatal stops /c, ɟ/ would be perceived as tokens of alveolar stops /t, d/ much in the same way that allophones of /t, d/ in English are perceived as belonging to the phonemic category of alveolar stops. As learning is error-driven, without accurate perception of the novel phonemic contrast new segmental representation cannot be established.

As I argued in Chapter Four, the distinguishing feature between Czech alveolar /t, d/ and palatal /c, ɟ/ is the monovalent feature [posterior] as a dependent of the Coronal Place node for the palatal segments /c, ɟ/. Moreover, while English does not have palatal stops in its phonemic inventory and therefore requires only a bare Place node in its representation of coronal stops (without further elaboration of the Coronal node by the feature [posterior]), it *does* require the feature [posterior] to distinguish between the front, or anterior, fricatives /s, z/ and back, or posterior, fricatives, /ʃ, ʒ/. Thus, English speakers *should* be able to perceptually contrast alveolar /t, d/ and palatal /c, ɟ/ since they manipulate the distinguishing feature for another phonemic contrast in their inventory. This gives us the hypothesis:

Hypothesis: Since English *does* utilise the feature [posterior] in its inventory for some phonemic contrast, then English speakers *should* be able to perceive the distinction between Czech alveolar and palatal segments. They will subsequently be able to establish a new segmental representation for palatal segments.

1.1 Methodology

1.1.1 Subjects

Eight adults self-reporting normal hearing participated in this experiment. The experimental group consisted of six adult North-American English speaking learners of Czech. All six subjects were living in Prague, Czech Republic and were enrolled in Czech language courses. Despite living in a Czech-speaking environment, all subjects reported that their main language of communication at work and socially was primarily English; Czech was used minimally. The subjects were between 25 and 40 years of age. There was a wide range in length of exposure to the Czech language, ranging from three months to ten years. Ages and length of exposure to Czech for each subject is indicated in (1). (Note that length of exposure is marked from the time of first arrival in the Czech Republic and not from the start-of-study date; moreover, some subjects had returned to the United States and back to Prague at least once.)

- (1)
- | | |
|-------------------|--|
| Subject 1: | ML, age 26 |
| | Length of exposure to Czech: 3 months |
| Subject 2: | JD, age 37 |
| | Length of exposure to Czech: 5 months |
| Subject 3: | AD, age 31 |
| | Length of exposure to Czech: 11 months |
| Subject 4: | SW, age 31 |
| | Length of exposure to Czech: 11 months |
| Subject 5: | JA, age 25 |
| | Length of exposure to Czech: 1 year |
| Subject 6: | RK, age 40 |
| | Length of exposure to Czech: 10 years |

There was also a control group consisting of one female and one male native speaker of the Prague dialect of the Czech language. The ages of the two control subjects were 26 and 28 years, respectively.

1.1.2 Stimuli

The stimuli consisted of 100 individual mono- and bi-syllabic words in which one of the four Czech test segments, i.e. voiceless alveolar /t/, voiced alveolar /d/, voiceless palatal /c/ or voiced palatal /j/, occurred. Test segments appeared either word-initially, word-medially, or word-finally, with the exception that the two voiced segments /d/ and /j/ did not appear word-finally. As voiced consonants devoice word finally in Czech, the contrast between /t ~ d/ and /c ~ j/ conflates at the surface level so it was impossible to include tokens of word-final /d/ and /j/. The breakdown of stimuli by word position of the test segment is shown in (2):

- (2) (a) 40 tokens with a word-initial stop: 10 each of word-initial [t], [c], [d], [j]
- (b) 40 tokens with a word-medial stop: 10 each of word-medial [t], [c], [d], [j]
- (c) 20 tokens with a word-final stop: 10 each of word-final [t], [c]

In order to examine the effect, if any, of following vowels (or preceding vowels in the case of word-final test tokens), vowel quality was varied with equal numbers of the vowels /a/, /ɛ/, /ɪ/, /o/, and /u/ following (or preceding) the test segments. For consistency, an effort was made to use only short vowels; however, in some instances a lexical item containing a short vowel did not exist so a word with a long vowel was used. This should not affect the task as English does not have contrasting vowel quantity. Thus, each of the four stop segments would have two tokens of each following vowel for word-initial and word-medial stimuli, and two tokens of each preceding vowel in the case of word-final tokens. The breakdown according to word-position as well as vowel quality is given in (3):

- (3) (a) 40 tokens with a word-initial stop
 = 10 tokens each of word-initial [t], [c], [d], [ʃ]
 = 2 tokens each of following vowel [a], [ɛ], [ɪ], [o], [u]
- (b) 40 tokens with a word-medial stop
 = 10 tokens each of word-medial [t], [c], [d], [ʃ]
 = 2 tokens each of following vowel [a], [ɛ], [ɪ], [o], [u]
- (c) 20 tokens with a word-final stop
 = 10 tokens each of word-final [t], [c]
 = 2 tokens each with *preceding* vowel [a], [ɛ], [ɪ], [o], [u]

To illustrate, a single token of each type of stimuli for the segment /t/ is given in (4).
 (The complete list of test stimuli in orthographic and phonetic forms can be found in Appendix One.)

(4)

	[a]	[ɛ]	[ɪ]	[o]	[u]
Word-initial [t]	taky	tebe	tyka	tolik	tužka
Word-medial [t]	jota	jetel	věty	beton	potuk
Word-final [t]	bělat	ojet	nosit	jekot	ochut

The stimuli were all tokens of natural Czech words. To ensure that the stimuli were identical for all subjects, the stimuli were recorded using a Sony TCD-100 Digital Audio Tape (DAT) Recorder and a Universal ECM-MS908C Stereo Microphone by a female 26-year old, University-educated native speaker of the Prague dialect of Czech. A timer was used during recording in order to leave a 10-second break between each token.

Before the experiment proper began, the experimenter prepared the subjects for the task with an overview of the four test segments as well as a five-token practice session. The overview consisted of a brief presentation of the four segments /t/, /d/, /c/, and /j/ in which each sound was produced in isolation and then in word-initial position. This was repeated three times for each segment, for a total of six repetitions of each segment. Following the overview presentation of sounds, subjects were given a practice run consisting of five tokens to familiarise them with the procedure and symbols. Subjects were told that each item they would hear contained one of the four target sounds, and that the sound could come at the beginning, middle, or end of the word. The subjects were told that they would be required to fill in blanks in the written words indicating which of the sounds they had heard, using the symbols shown in (5) for each sound:

(5)	If you hear:	Then write:
	[t]	t
	[c]	t'
	[d]	d
	[j]	d'

The subjects were asked if they understood the test procedure: if they indicated that they understood, then we proceeded to the practice stimuli. Subjects were given a practice sheet on which to mark the perceived sounds. The practice sheet is shown in (6):

(6)	Practice:
	a. u__ek
	b. __ole
	c. chres__
	d. __uka
	e. o__iv

Following the practice, we proceeded to the experiment proper regardless of the correctness or incorrectness of the subject's responses on the practice tokens. Subjects were not given feedback on their responses. For the experiment, the subjects were presented with a sheet of paper with a list of words numbered from 1-100. As in the practice, a blank was left in each word indicating where the subject was to mark the perceived segment. (A sample of the test sheet can be found in Appendix One.) The words were in random order. The aural stimuli for both the five practice tokens and the test proper were presented over headphones attached to the DAT recorder at a comfortable listening level as determined by the subject. After each 25 words, the subjects were given a 30 second break to avoid test fatigue.

1.2 Results

For the Forced Choice Phoneme Selection task, tokens were marked on place of articulation and not voicing quality so that if a subject marked alveolar 't' rather than alveolar 'd', it was considered to be correct. As alveolar stops occur in English, performance on alveolar tokens was not counted in the tabulation of percentages as accurate perception on these items would be consistent with the subjects' L1. Only one subject perceived a token of an alveolar stops as a palatal. Results show that the subjects could perceive the Czech alveolar /t, d/ versus palatal /c, ʃ/ distinction at greater than chance levels, and results varied with length of exposure to Czech. (Individual and group results for the FCPS task can be found in Appendix Two.) The two Czech controls both scored 100%. Table 5.1 illustrates the percentage of correct tokens for each subject, by word position:

Table 5.1 Percentage of Palatal Stop Tokens Perceived Correctly

	Subject / Length of Exposure					
Word Position	Subject 1 ML 3 months	Subject 2 JD 5 months	Subject 3 AD 11 months	Subject 4 SW 11 months	Subject 5 JA 1 year	Subject 6 RK 10 years
<i>Initial</i>	70	90	80	85	80	95
<i>Medial</i>	70	70	80	90	85	90
<i>Final</i>	20	30	50	70	70	80

Let's look first at the overall performance of the group. Overall, all six subjects had the greatest perceptual acuity for tokens with word-initial palatal stops: palatals were perceived as alveolar /t, d/ less frequently in word-initial position. Percentages correct for word-initial palatal tokens ranged from a low of 70% for Subject 1, ML, a beginner with only three months of exposure to Czech, to a near perfect score of 95% (only one token perceived incorrectly) for RK with 10 years of exposure to Czech. Higher performance on word-initial tokens can be seen as a result of greater saliency for word-initial consonants, as place of articulation cues are provided by the change from consonant to the following vowel.

Conversely, word-final palatal stops were correctly perceived as palatals at a much lower rate. This can be expected as acoustic cues for place are less salient word-finally without recourse to place of articulation cues provided by transition effects of the shift from consonant to vowel. The two subjects with the least amount of exposure to Czech were at lower than chance levels on word-final palatal stops: Subject 1, ML, perceived 20% of the tokens containing word-final voiceless palatal [c] correctly while Subject 2, JD, has a slightly higher score with 30% of the word-final voiceless palatal [c] perceived correctly. However, these two scores on word-final tokens were the only scores below chance of the entire FCPS task. A third subject, AD, correctly perceived the word-final palatal stops at chance levels of 50%. These three scores on word-final [c] are much lower than the overall scores for all tokens. Subjects 4, 5 and 6 with relatively more exposure to Czech than Subjects 1, 2 and 3 perceived word-final palatal stops at

greater than chance levels. The highest percentage of correctly perceived word-final stops was 80% by RK, the subject with the most Czech experience. With respect to following vowel quality, overall all subjects performed significantly worse on palatal tokens with a following high front vowel unrounded /ɪ/.

We now take a look at individual performances. In general, individual performances corresponded to the overall results in that all six subjects showed greatest perceptual acuity for word-initial palatal /c, j/, and worst acuity for word-final /c, j/. I will go over the subjects in order of increasing length of exposure to Czech.

Subject 1, ML, had three months of Czech exposure and misperceived a total of twenty palatal tokens out of a possible fifty. As did all subjects, she performed significantly worse on tokens with a following high front unrounded /ɪ/; however, in ML's case the difference between performance on tokens with a high front vowel /ɪ/ and other vowel qualities was markedly different. ML made two errors on following vowel /a/ and a single error on each of the following vowels /ɛ/, /o/ and /u/; but she had *seven* errors on tokens with a high front unrounded /ɪ/. ML also had difficulty perceiving word-final palatal stops, with eight out of ten word-final palatals perceived incorrectly as alveolar stops. As noted above, lower perceptual acuity for word-final consonants is expected as word-final consonants lack place of articulation cues provided by the transition from consonant to vowel. ML's performance on word-initial and word-medial palatal stops was at above chance levels, with six out of twenty errors for each position. In terms of perceptual differences between voiceless and voiced palatal stops¹, there was no difference in perceptual acuity between voiced palatal /j/ versus voiceless palatal /c/: ML made six errors on each type.

The second subject, JD, had 5 months of exposure to Czech. He misperceived a total of 17 palatal tokens out of the total possible fifty. JD's perceptual acuity for the Czech stop contrast was similar to the other five subjects in that his score for word-initial

¹ Error calculation on the basis of voicing quality of the palatal (i.e. [c] ~ [j]) is calculated only for word-initial and word-medial tokens. Word-final tokens were not included in the calculation as stops devoice word-finally in Czech: including only word-final voiceless stops would skew the results.

palatal tokens was higher than his score for word-final tokens. As with the other subjects, he made four errors on the most commonly misperceived group of palatal tokens with a following high front vowel unrounded /ɪ/ and a single error on each of the tokens with following back vowels /o/ and /u/. With respect to differences in performance on voiceless /c/ versus voiced /j/, JD had greater perceptual acuity for voiced palatal /j/ in making three errors on voiced exemplars compared to five on voiceless tokens.

JD's results were idiosyncratic in a number of ways. JD was the only subject to incorrectly perceive *alveolar* tokens as *palatal* stops in two instances: (i), word-medially before a high front vowel unrounded /ɪ/ as in token number 47, 'kudy' [kudɪ]; and (ii), before a mid-front vowel /ɛ/ as in token number 27, 'někdo' [ɲɛɡdo]. JD was also unusual in that his perceptual acuity for word-medial palatal tokens was substantially worse than his acuity for word-initial palatals. While most subjects had equal or near-equal scores on word-initial and word-medial tokens of /c, j/, JD incorrectly perceived two out of twenty word-initial palatals incorrectly in word-initial position and six out of twenty word-medial palatals. In terms of vowel quality, when the vowel following the palatal stop was a low-front vowel /a/, JD incorrectly perceived palatal /c, j/ as alveolar /t, d/ at a higher rate than did the other subjects: while four of the six subjects correctly perceived *all* instances of palatals with a following /a/, JD incorrectly perceived three palatals tokens with a following /a/.

Subjects 3, 4, 5 and 6 had similar results in that they correctly perceived *all* palatal tokens with following front vowels /a/ and /ɛ/. Subject 3, AD, came to the task with 11 months exposure to Czech. She made a total of 13 errors out of the fifty exemplars containing palatals, or 26% incorrect. AD made an equal amount of errors on word-initial and word-medial palatals stops with four out of 20 palatals, or 20%, perceived incorrectly. She made substantially more errors on tokens with a following high front vowel unrounded /ɪ/ than on other vowel qualities, incorrectly perceiving five palatals as alveolars with a following /ɪ/ as compared to a single error with a following mid-back vowel /o/ and two errors with a following high-back vowel /u/. In terms of

perceptual differences on voiceless versus voiced palatals, AD had the opposite tendency from Subject 2, JD, in that she misperceived three voiceless palatals /c/ and five voiced palatals /j/.

The fourth subject, SW, had the same length of exposure as AD at eleven months yet performed significantly better than AD on the perception task: she misperceived only eight out of the fifty palatal tokens, or 16%, as compared to the 13 tokens, or 26%, incorrectly perceived by AD. Moreover, SW provided the single exception to the general finding that perceptual ability to distinguish Czech alveolar /t, d/ and palatal /c, j/ increased as a function of exposure to Czech in that she made *fewer* perceptual errors than did Subject 5, JA, with one year of experience. Like Subject 3, AD, she made no perceptual errors on tokens with following front vowels /a/ and /ε/. However, she incorrectly perceived three palatals with a following high front vowel unrounded /ɪ/ as compared to the five errors made by AD, and a single error on each of the palatal groups with a following back vowels /o/ and /u/. Her performance on word-initial and word-medial palatals was similar with three out of twenty errors on word-initial palatal stops and two out of twenty errors on word-medial palatals. SW's performance was not significantly influenced by voicing quality of the palatal with three errors on voiceless palatal /c/ tokens compared to two on voiced palatal /j/ tokens.

JA, the fifth subject with one year of experience, had a total of ten errors on the palatal tokens out of a possible fifty, or 20% incorrect. As with Subjects 3 and 4, JA made *no* perceptual errors on tokens with following front vowels /a/ and /ε/. She had similar performance on the three other preceding vowel qualities with three errors on palatal stop tokens preceding the high front vowel unrounded /ɪ/; and two errors on each of the following vowels back /o/ and /u/. JA showed a marked difference in perceptual acuity with respect to voicing quality of the palatal: she made two errors on exemplars with a voiceless palatal /c/ and five with a voiced palatal /j/.

The sixth subject, RK, had substantially more exposure to the Czech language than any of the other five subjects and had the best overall score. The number of misperception errors was correspondingly lower: he had a total of five misperceived

palatal tokens out of the fifty exemplars, or 10%. Out of the total errors, two of the five misperceptions were on word-final tokens on which place of articulation cues are less salient: he incorrectly perceived the word-final palatal /c/ in tokens number 53 “*havět*” and number 80 “*plet*”. RK made two errors on word-medial tokens and one on word-initial. In terms of vowel quality, RK incorrectly perceived two tokens with following high front vowel /i/ and one with following high back vowel /u/.

1.3 Discussion

The overall results of the Forced Choice Phoneme Selection task showed that English speakers can perceive the distinction between Czech alveolar /t, d/ and palatal /c, j/ at greater than chance levels. These perception results supported the hypothesis proposed in this thesis. Thus, if Brown’s model is correct, English speakers should be subsequently able to establish new segmental representations for the palatal segments /c, j/. That is, they will be able to contrastively perceive and produce alveolar /t, d/ versus palatal /c, j/. In the next section, we will see if the production tasks support this hypothesis. Before turning to the production experiments, however, we will discuss the findings and implications of the FCPS task in more detail.

An interesting result was obtained in the perception experiment in that the subjects’ scores on the FCPS task were related to their length of exposure to Czech: the more exposure, the higher the score. These results are encouraging: recall from Chapter Three, Section 3.3 that Brown’s (1998) Japanese subjects showed no increase in perceptual acuity on the English lateral approximant /l/ versus central approximant /r/ contrast with increased exposure to English which Brown linked to the absence of the feature [Coronal] in the Japanese feature inventory.

As the speakers I tested were able to phonemically contrast alveolar /t, d/ and palatal /c, j/ with greater acuity with increased exposure to Czech, they may be able to improve on their productive ability with time. Their experience in perceiving the phonemic contrasts between the alveolar fricatives /s, z/ and the alveo-palatal fricatives

/ʃ, ʒ/ along the acoustic dimension defined by the feature [posterior] enables them to discriminate non-native contrasts along that same dimension.

Recall also from Chapter Three, Section 3.4 that Matthews (1997) measured the effects of pronunciation training on non-native contrasts in which the Japanese subjects showed no increase in perceptual acuity on the English /l/ versus /r/ contrast following pronunciation training, while they showed significant improvement on the non-native contrasts such as /b ~ v/ and /s ~ θ/. Matthews argued that the segmental contrasts showing improvement were distinguished along dimensions corresponding to features in the Japanese L1 feature geometry, while the segmental contrast showing no improvement after training was characterised by a feature completely absent from Japanese. Thus, if the English speakers I tested can show increased acuity for the segmental contrast as a result of increased exposure, it is predicted that their production will improve.

Carroll (1999) points out that L2 perceptual abilities to detect properties of the L2 signal appear to vary as a function of the lexicon. She notes that English speaking learners of French, Greek or Spanish will initially 'hear' word-initial voiceless stops as voiced stops so that the Greek *‘pino krazī’* "I drink wine" will be heard as either *‘bino grazi’* or *‘bino krazī’* due to the long-lag VOT of English. Once the learner realises that the word *‘bino’* would not make sense in that context he or she will realise that a slip of the ear has occurred.

In sum: the six English-speaking learners of Czech tested in the Forced Choice Phoneme selection task were able to perceive the Czech alveolar /t, d/ versus palatal /c, ʃ/ distinction at greater than chance levels. Word-position of the palatal played a role in perceptual acuity, with word-initial palatal /c, ʃ/ tokens perceived correctly more frequently than word-final /c, ʃ/ due to place of articulation cues provided by transition effects from consonant to following vowel. Under Brown's phonological model of L1 interference assumed in this thesis, learners of a second language should be able to perceive a novel segmental contrast if that contrast is distinguished by a feature present elsewhere in the native inventory, as the presence of the contrasting feature in the L1 will allow them to perceive any contrasts along that dimension. As the subjects were able to

perceptually contrast Czech alveolar versus palatal stops, the results of this experiment thus appear to support the hypothesis that English does manipulate the feature [posterior] in its inventory. Because English does manipulate the contrasting feature, English-speaking learners of Czech should eventually be able to productively contrast Czech alveolar /t, d/ versus palatal /c, ʃ/. In the next section we turn to production tasks to see if these English-speaking learners of Czech can contrastively produce the Czech alveolar /t, d/ versus palatal /c, ʃ/ distinction.

2.0 EXPERIMENT TWO: PRODUCTION TASKS

As noted in the introduction, in order to avoid over- or under-estimating adult L2 learners' phonological or mental representations for novel segmental contrasts, we require two types of tasks to determine whether or not second language learners have acquired a non-native phonemic contrast: production and perception tasks. In Experiment 1, I established that English speaking learners of Czech can perceptually contrast alveolar stops /t, d/ with palatal stops /c, ʃ/ at greater than chance levels, and that perceptual acuity increased with length of exposure to Czech. Now it is time to turn to production evidence for the Czech alveolar /t, d/ versus palatal /c, ʃ/ segmental contrast. Research has shown that both speakers who can perceive the novel phonemic contrast and those who cannot are able to produce the novel segment.

2.1 Methodology

2.1.1 Subjects

A random four adult Czech learners from Experiment 1 participated in Experiment 2, including both the subjects with the least and the most amount of exposure to Czech. Two subjects, AD and SW, each with eleven months of experience did not participate. Subjects, ages, and length of exposure are given in (7):

- (7) **Subject 1:** ML, age 26
Length of exposure to Czech: 3 months
- Subject 2:** JD, age 37
Length of exposure to Czech: 5 months
- Subject 3:** JA, age 25
Length of exposure to Czech: 1 year
- Subject 4:** RK, age 40
Length of exposure to Czech: 10 years

2.1.2 *Stimuli*

To elicit speech data for the production experiment, subjects were recorded in several situations: reading a list of 15 sentences; responding to questions in casual conversation; and spontaneous speech whenever possible. The stimuli for the sentence reading task can be found in Appendix Three. The stimuli sentences varied in length: each contained a minimum of three tokens of test segments /t/, /d/, /c/, or /y/, up to a maximum of seven. The token-containing words consisted of both high and low frequency words. I will first discuss tokens obtained in free speech before moving on to the Sentence Reading task.

2.2 *Results*

2.2.1 *Free Production*

The samples for the free production task were obtained by either questioning in conversation or random samples of spontaneous speech recorded when possible. Over the course of the data elicitation, the four subjects produced a variety of words containing tokens of alveolar and palatal stops. In the interests of clarity and comparison, I present only those words that were produced by all four speakers. Results are given in Table 5.2 below:

Table 5.2 English Speakers' Production of Palatal Stops in Free Speech

Orthographic Form	Phonetic Form	Gloss	Subject / Length of Exposure			
			Subj.1	Subj. 2	Subj.3	Subj.4
			ML 3 mo.	JD 5 mo.	JA 1 yr	RK 10 yrs
1. delám	[jɛla:m]	'do' (1.p.s)	[j]	[d]	[dj]	[d]
2. rodinu	[roʝɪnu]	'family'	[d]	[d]	[d]	[d]
3. vidět	[viʝɛt]	'to see'	[j]	[d]	[dj]	[d]
4. díky	[ji:ki]	'thanks'	[j]	[d]	[d]	[d]
5. věděla	[vɛʝɛla]	'knew' (3.p.s.)	[d]	[d]	[dj]	[dj]
6. divadlo	[ɟɪvadlo]	'theatre'	[j]	[d]	[d]	[dj]
7. děti	[ʝɛci]	'kids'	[d / t]	[d / t]	[dj / t]	[dj / t]
8. ti	[ci]	'you (dat.)'	[t]	[t]	[t]	[t]
9. ještě	[jɛʃɕɛ]	'still'	[t]	[t]	[tj]	[t]
10. teď	[tɛɕ]	'now'	[t]	[t]	[t]	[t]
11. město	[mjɛsɕɛ]	'(in) town'	[tɛʝɛ]	[t]	[tj]	[tj]

Let's look at the group patterns first. A crucial finding was that no subject produced native-sounding palatal stops. Both inter- and intra-speaker variability was common, with all subjects producing more than one substitution in place of palatal stops /c, ɟ/. There were three types of substitution for palatal stops. The three patterns of substitution are given in (8) in descending order of frequency of occurrence across subjects:

- (8) Substitution for Czech Palatal /c, ɟ/ in Free Production
- (i) Alveolar [t, d]
 - (ii) Alveolar stop plus palatal glide sequence [tj, dj]
 - (iii) Palatal glide [j]

The most common pattern was alveolar substitution, with all subjects substituting alveolar [t, d] for palatal [c, ɟ] in at least some instances. The second most common pattern was alveolar stop plus palatal glide sequences [tj] and [dj] as produced by Subjects 3 and 4. The third pattern of substituting a palatal glide [j] for the voiced palatal stop [ɟ] was produced by only one speaker, ML.

Now let's look at individual performances. Both inter- and intra- speaker variation in production were common, yet the variation showed consistent patterns. For each subject, I will present individual patterns of production in the natural speech task, followed by a comparison with performance on the tokens elicited on the reading task discussed in Section 2.2.1. Individual performances are discussed in ascending order of length of exposure to Czech.

ML, the subject with the least amount of exposure to Czech at 3 months, systematically produced different segments for the voiced and voiceless palatals. For the voiceless Czech palatal, ML had a single substitution: she produced voiceless alveolar stops [t] in place of the voiceless palatal stop [c] in all instances. There were no instances of palatal glides replacing voiceless palatal stop which I attribute to the fact that English does not have a phonemic voiceless palatal glide. Examples of ML's production for voiceless palatals are presented in (9). (For this example and all the following examples in this section on free production, the number in square brackets directly to the left indicates the token number on the list in Table 5.2):

(9)	orthographic form	Czech production	ML's production
a.	[8.] ti	[ci]	[ti]
b.	[9.] ještě	[jɛʃcɛ]	[jɛʃtɛ]
c.	[10.] ted'	[tɛc]	[tɛt]

However, for the voiced palatal stop [j] ML alternated between production of the palatal glide [j] in place of a stop of any kind, and the voiced alveolar stop [d]. Examples of ML's production of the palatal glide [j] in words containing the voiced palatal stop [j] are given in (10a-f): samples of her production of voiced alveolar [d] in place of the voiced palatal [j] are in (10d-f)

(10)	orthographic form	Czech production	ML's production
a.	[1.] dělám	[jɛlam]	[jɛlam]
b.	[3.] vidět	[viɟet]	[viɟet]
c.	[4.] díky	[ji:kɪ]	[jiki]
d.	[2.] rodinu	[rojɪnu]	[rodinu]
e.	[5.] věděla	[vjɛɟɛla]	[vjɛdɛla]
f.	[7.] děti	[jɛci]	[dɛti]

The two patterns of production for Czech palatal segments /c, j/ indicate that ML is beginning to productively contrast alveolar versus palatal stops although her production is not nativelike. However, this productive contrast is not consistent as she is still producing alveolar /t, d/ in place of palatal /c, j/ in some instances.

At five months of exposure to Czech, Subject 2 JD did not produce any palatal glides or alveolar stop plus glide sequences, but instead substituted alveolar stops for both voiced and voiceless palatal stops in all instances. In doing so, JD fully conflated the phonemic distinction between alveolar /t, d/ and palatal /c, j/. Thus, JD produced alveolar [t] for both alveolar [t] and palatal [c] as shown in figure (11a-b); he produced voiced alveolar [d] for both alveolar [d] and palatal [j] as shown in (11c-d):

(11)	orthographic form	Czech production	JD's production
a.	ty	[tɪ]	[tɪ]
b.	[8.] ti	[ci]	[ti]
c.	jdu	[du]	[du]
d.	[1.] dělám	[jɛla:m]	[dɛlam]

Thus, while JD's results on the FCPS task showed his perceptual acuity to be at greater than chance levels, he does not appear to productively contrast the phonemes in production.

The fourth subject JA, at one year of exposure to Czech, produced both alveolar stops /t, d/ as well as alveolar plus glide sequences /tj, dj/ for both voiced and voiceless palatal segments. However, of all four subjects JA was the most consistent in her production of alveolar plus palatal glide sequences in place of palatal stops: that is, JA conflated the distinction between alveolar /t, d/ and palatal /c, j/ less often than did the other subjects. Vowel quality influenced production: JA consistently produced an alveolar stop plus glide [tj, dj] in place of both voiceless and voiced palatal stops when the following vowel was a mid front unrounded /ɛ/. JA's production for voiced palatal stops before the mid front unrounded vowel /ɛ/ is shown in figure (12a-d): her production for voiceless palatal stops before the mid front unrounded vowel /ɛ/ is shown in (12e-f):

(12)	orthographic form	Czech production	JA's production
a.	[1.] dělám	[jɛla:m]	[djɛlam]
b.	[3.] vidět	[viɛt]	[vidɛt]
c.	[5.] věděla	[vjɛjɛla]	[vjɛdjɛla]
d.	[7.] děti	[jɛci]	[djɛti]

e.	[9.]	ještě	[jɛʃcɛ]	[jɛʃtjɛ]
f.	[11.]	měště	[mjɛscɛ]	[mjɛstjɛ]

In front of all other vowels, JA produced a plain alveolar [t, d] rather than a palatal stop or alveolar plus glide sequence for both voiceless and voiced palatal stops. This is presented in (13):

(13)	orthographic form	Czech production	JA's production
a.	[2.]	rodinu	[rojmu]
b.	[8.]	ti	[tɪ]
c.	[6.]	divadlo	[ɟɪvadlo]

The fourth subject, RK, also varied in production between alveolar and alveolar plus glide sequences. In place of Czech palatal stops /c, ʃ/, RK often produced alveolar plus palatal glide sequences [tj, dj] as shown in (14):

(14)	orthographic form	Czech production	RK's production
a.	[5.]	věděla	[vjɛjɛla]
b.	[6.]	divadlo	[ɟɪvadlo]
c.	[7.]	děti	[ɟɛtɪ]
d.	[11.]	měště	[mjɛstjɛ]

However, RK's alternation between alveolar stop and alveolar stop plus glide sequences did not appear to be influenced by the following vowel quality as we saw

earlier with Subject 3, JA. As discussed above, JA consistently produced an alveolar plus glide sequence when the following vowel was a mid front unrounded vowel /ɛ/ and an alveolar stop at other times; RK was not as consistent in that he produced both alveolar stops and alveolar plus glide sequences when the following vowel was a mid front /ɛ/. Exemplars of RK's production on words with a palatal followed by a mid front vowel /ɛ/ are given in (15); voiced and voiceless alveolar plus glide sequences are given in (15a-b); voiced and voiceless alveolar stops are given in (15c-d):

(15)	orthographic form	Czech production	RK's production
a.	[5.] věděla	[vɛɛjɛla]	[vjɛdjɛla]
b.	[11.] městě	[mjɛscɛ]	[mjɛstjɛ]
c.	[1.] dělám	[ɟɛla:m]	[dɛlam]
d.	[9.] ještě	[jɛʃcɛ]	[jɛʃtɛ]

2.2.2 Sentence Reading Task

For the sentence reading task, tokens were marked on place of articulation and not voicing quality so that if a subject said alveolar [d] rather than alveolar [t], it was considered to be correct. As with the perception tasks, performance on alveolar tokens was not counted in the tabulation of percentages as accurate perception on these items would be consistent with the subjects' L1. The two Czech controls both scored 100%. Results are in Appendix Three. The stimuli were initially transcribed by myself, with further consultation with two native speakers of Czech.

Turning first to the group patterns, the most noticeable finding was that no subject produced a native-sounding Czech palatal stop /c/ or /j/. The closest production to a Czech palatal stop was a sequence of an alveolar stop followed by a palatal glide. Palatal stops [c] and [j] were consistently replaced by one of two patterns: (i), alveolar [t, d]; or (ii), alveolar stop plus palatal glide sequences [tj, dj]. The first pattern whereby palatal

[c, ʝ] were replaced with alveolar [t, d] was significantly more frequent than the second pattern of alveolar plus glide sequences. The determining factor as to whether an alveolar stop or an alveolar plus glide was substituted appears to be lexical knowledge: palatal stops in common or high frequency words (for example, "děti" [jetɪ] ('*children*')) were replaced with alveolar plus glide sequences, while the palatal stops in less common words (for example, "choť" [xoc] ('*mother-in-law*')) were replaced with alveolar stops. In some instances, all palatal tokens in an entire sentence were replaced with alveolar stops by all subjects.

As we saw in Chapter 4, Section 1.2.1, the distinction between alveolar /t, d/ and palatal /c, ʝ/ is marked directly in the orthography in Czech; however, the cues can be confusing for a non-native speaker. This may be attributed partially to the fact that the orthographic cue indicating palatal /c, ʝ/ varies with respect to the following vowel so that palatal segments are marked in three different ways. If learners of Czech are disregarding the orthographic cues to place of articulation, then they cannot be expected to produce the palatal stop in unfamiliar lexical items in a reading task. As we will see shortly in the discussion of individual results, one subject, JA, did appear to be sensitive to the orthographic cue for palatal stop on tokens with a following 'e', i.e. 'ě'.

Now let's look at individual performances. Inter-speaker variation was common, yet within this variability there was a consistent pattern. The four subjects fell into two categories in terms of production for palatal stops: Group 1 conflated the phonemic distinction between Czech alveolar /t, d/ and palatal /c, ʝ/ while Group 2 maintained a phonemic distinction between Czech alveolar and palatal stops in some (but not all) cases, by producing both alveolar [t, d] and alveolar stop plus glide sequences [tj, dj] for the Czech palatal stops /c, ʝ/. Production for the two groups is given in (16):

(16) Palatal Production in the Sentence Reading task

Group One: Alveolar /t, d/ produced as: [t, d]

Palatal /c, j/ produced as: [t, d]

Group Two: Alveolar /t, d/ produced as: [t, d]

Palatal /c, j/ produced as: (i) [t, d]

(ii) [tj, dj]

Groups 1 and 2 were divided by amount of exposure to Czech. Group 1 consisted of the two early beginner subjects with the least amount of exposure to Czech: ML, with 3 months of exposure to Czech, and JD, with 5 months of exposure to Czech. Both subjects in this group produced alveolar [t, d] for palatal [c, j] in all instances in the reading task, thereby completely conflating the phonemic distinction between alveolar and palatal stops.

Group 2 consisted of the two subjects with more exposure to Czech: JA, with one year experience with Czech, and RK, with 10 years of exposure. The two subjects in the second group had two tendencies when faced with Czech palatal stops: they produced either an alveolar stop [t, d] or an alveolar stop plus glide sequences [tj, dj]. However, within this alternation was a consistent pattern. Both JA and RK appeared to produce an alveolar stop when the word containing a palatal segment was unfamiliar, and an alveolar stop plus glide sequence when the palatal-containing word was known to them. Moreover, JA consistently produced an alveolar stop plus glide sequence in place of a palatal stop when the palatal stop was followed by an *ě* in the orthography. As we saw in Chapter Four, Section 1.2.1 and discussed above, when the following vowel is [ɛ] the palatal place of articulation is more obviously indicated in the orthography by means of the ‘háček’ on the vowel, i.e. *ě*. The ‘háček’ above the ‘e’ appears to be a more salient cue than the alternating i/y or apostrophe used to indicate a palatal stop before other vowels.

2.3 Discussion

The results of the production experiments showed that English speaking learners of Czech can productively contrast alveolar /t, d/ versus palatal /c, ɟ/. although their productions of Czech palatal stops /c, ɟ/ are not fully nativelike. Subjects 3 and 4 who were able to contrast alveolar versus palatal stops produced alveolar plus glide sequences /tj, dj/ in place of simple palatal segments /c, ɟ/. The question arises as to the distinction between /tj/ and a palatal /c/, which may not appear to be a significant difference to a non-native speaker of Czech. Native speaker informants responded that the alveolar plus glide sequence sounded much longer: they were readily able to distinguish between alveolar /t, d/ versus palatal /c, ɟ/ versus the alveolar-glide sequence /tj, dj/ for a three-way contrast.

The two native speakers who performed as control subjects noted that while the alveolar plus glide sequence [tj, dj] was understood as a palatal stop due to lexical cues, they did not consider it to sound "Czech-like". They also noted that both the alveolar plus glide and the alveolar substituted for the palatal were preferable to substituting a palatal glide [j] which they considered to be unrecognizable as an exemplar of a palatal stop. Without contextual cues, the lexical item would be unrecognizable if produced with a palatal glide rather than a voiced palatal stop. Thus, while three of the four English speaking subjects could phonemically contrast alveolar versus palatal stops, the challenge is for them to move their articulations back to the palate rather than producing the substantially more front alveolar plus glide sequences.

Because the production experiment involved both a reading task as well as production in free speech, I was able to compare subjects performance across production tasks. In the sentence reading task, knowledge of orthographic conventions played an important role in subjects' productions. As we saw in Chapter Four, Section 1.2.1, both alveolar /t, d/ and palatal /c, ɟ/ sounds use the symbols 't' and 'd' in orthography with the palatal being distinguished orthographically from the alveolar by means of a diacritic. Thus, the distinction between alveolar /t, d/ and palatal /c, ɟ/ is marked directly in the orthography in Czech; however, comparing the results of the Sentence Reading task with results obtained in free speech, these orthographic cues do not appear to be especially

salient to a non-native speaker of Czech. As I noted above, the problem may be that the orthographic cue are not consistent across segments, but rather vary with respect to the following vowel so that palatal segments are distinguished from alveolar in three different ways: with a following apostrophe, a change in vowel symbol, and a Czech 'háček'. If the subjects are disregarding the cues, then they cannot be expected to produce the palatal stop in unfamiliar words in a reading task.

Another confounding variable may be that the Czech diacritic 'háček' is used orthographically for purposes other than indicating a palatal stop before an 'e'. First, the 'háček' is used above the letters 's', 'z' and 'c' (i.e. 'š', 'ž', and 'č') to indicate that they are pronounced [ʃ], [ʒ], and [tʃ], respectively; it is also used above the 'r' (i.e. 'ř') to indicate the trill fricative [r]. Secondly, and more importantly for our purposes, the 'háček' is also used following the bilabial and labiodental segments [b], [p], [v], [f] and [m] to indicate that the stop or fricative is followed by a glide; that is, the letters 'b', 'p', 'v', 'f' and 'm' followed by an 'ě' are pronounced [bj], [pj], [vj], [fj] and [mj], respectively. Short (1993: 459) notes that the "ě after b, p, f, v denotes not palatalised labials (lost in the fifteenth century) but a fully developed palatal element [j]." Examples are given in (17):

(17)	Orthographic form	Phonetic form	Gloss
a.	bělat	[bjelat]	'to groan'
b.	pět	[pjɛt]	'five'
c.	město	[mjɛsto]	'city'
d.	věděla	[vjɛɟɛla]	'knew' (3.p.s)

The use of the 'háček' as a cue to the palatal glide following the stops [p], [b], [m] and fricatives [f] and [v] may shed some light on why the native English speaking subjects I tested produced palatal /c, ɟ/ as alveolar-glide sequences in many instances.

even after ten years of exposure to Czech as in the case of RK. In terms of frequency, the majority of cases where an 'e' with a 'háček' over it occurred were for instances other than palatal stops /c, ʃ/. Moreover, there are many word pairs that differ orthographically only by the consonant preceding the vowel so that they appear to be minimal pairs.

Examples are shown in (18):

(18)	Orthographic form	Phonetic form	Gloss
a.	město	[mjesto]	'city'
b.	těsto	[često]	'dough'

Recall that the four subjects who participated in the production tasks fell into two categories with respect to production for palatal stops in the Sentence Reading task:

Group 1 fully collapsed the phonemic distinction between Czech alveolar and palatal stops while Group 2 maintained a phonemic distinction between Czech alveolar and palatal stops in some (but not all) cases, by producing both alveolar [t, d] and alveolar stop plus glide sequences [tj, dj] for the Czech palatal stops /c, ʃ/.

Group 1 consisted of the two subjects with the least amount of exposure to Czech: ML, with 3 months of exposure to Czech, and JD, with 5 months of exposure to Czech. Both subjects in this group produced alveolar [t, d] for palatal [c, ʃ] in all instances in the reading task. Given that both subjects had very little exposure to Czech, it is not surprising that they did not pick up on the orthographic cues to the alveolar versus palatal distinction. Moreover, as much of their experience was in a conversational setting, they may not have recognized the written form of the lexical items they were familiar with only in spoken form. Several of the words containing palatal stops which appeared in the sentence reading task were also produced by ML and JD in the course of casual conversation. While JD fully collapsed the distinction between alveolar /t, d/ and palatal /c, ʃ/ in free production as well as in the sentence reading task, ML did show some variation in production of voiced palatal tokens in free speech, indicating an awareness

that alveolar /d/ and palatal /j/ are in contrast. For the voiced palatal stop [j], ML occasionally produced the voiced palatal glide [j] in place of the palatal stop [j]. For example, ML produced the palatal stop [j] in the word "děláám" [jɛlám] ('I do') with a voiced alveolar [d] in the Sentence Reading task while in the course of conversation she produced it as [jɛlám] with a voiced palatal glide [j]. However, for the voiceless palatal counterpart /c/, ML consistently produced only voiceless alveolar stops [t] in place of the voiceless palatal stop [c]. There were no instances of palatal glides replacing voiceless palatal stops which I attributed to the fact that English does not have a phonemic voiceless palatal glide.

Group Two consisted of the two subjects with relatively more exposure to Czech: JA, with one year experience with Czech, and RK, with 10 years of exposure. The two subjects in Group Two productively contrasted alveolar /t, d/ with palatal /c, j/. although their productions for /c, j/ were not 100% consistent or nativelike. Group Two had two tendencies in production when faced with Czech palatal stops: they substituted either an alveolar stop [t, d] or an alveolar stop plus glide sequences [tj, dj]. The alternation was consistent in that they produced an alveolar stop plus glide sequence when the palatal-containing word was known to them (as evidenced by their production of the same lexical items in natural speech) and an alveolar stop when the word containing a palatal segment was a low frequency word. JA also appeared to be aware of orthographic conventions to some extent in that she produced an alveolar stop plus glide sequence in place of a palatal stop in all instances in the reading task when the palatal stop was followed by an ě in the orthography; this occurred even in low-frequency words such as the old-fashioned word "děvče" [jɛvtʃɛ] ('lass').

3.0 STAGES OF ACQUISITION: CZECH PALATAL STOPS

The results of the Forced Choice Phoneme Selection task showed that the English speakers tested could perceptually contrast the Czech alveolar /t, d/ versus palatal /c, j/ at greater than chance levels. On this basis, a prediction was made that learners of Czech

could productively contrast these two phonemic pairs, rather than producing palatal stops as tokens of alveolars. I found that English speaking learners of Czech with some experience were able to productively contrast alveolar /t, d/ with palatal /c, j/, but that production varied as a result of experience and length of exposure to Czech. No speaker was able to produce a native sounding Czech palatal stop. However, within the variety there is a consistent pattern. The production of Czech palatal stops /c, j/ by the native English speakers I studied can be broken down into two stages:

Stage 1: Alveolar stop [t, d]

Conflating phonemic distinction between alveolar /t, d/ and palatal /c, j/

Both /t, d/ and /c, j/ funneled into alveolar category

Stage 2: Alveolar stop plus palatal glide sequence [tj, dj]

Phonemic contrast between alveolar /t, d/ and palatal /c, j/

Stage One is preceeded by a "pre-acquisition" stage, in which the voiced palatal stop /j/ is perceived and produced by native English speakers as a voiced palatal glide [j]. The substitution of [j] in place of /j/ is frequently produced by native speakers of English on initial exposure to the language. However, the voiceless palatal counterpart /c/ is generally produced as an alveolar [t]. I attribute this to the fact that English has a voiced palatal glide [j] in its phonemic inventory, but no voiceless phonemic counterpart.

Native-speaker informants reported that producing the glide [j] rather than a voiced palatal stop renders the word unintelligible. They noted that this type of substitution occurred when the English speaker had little or no prior knowledge of the Czech language and was repeating a word uttered by the native Czech speaker. The common example cited was the Czech word for "thank you", i.e. 'děkuju' [ɟɛkuju]. Thus, if they knew what the speaker was trying to produce (if, for example, the person had asked how to pronounce a word and was subsequently repeating it), then they could

recognize it as an instance of that particular lexical item. otherwise, substituting a palatal glide for a palatal stop renders the word unintelligible.

Subject One, ML is appears to be at this pre-acquisitional stage: she can perceptually contrast alveolar /t, d/ and palatal /c, ɟ/ at greater than chance levels in word-initial and word-medial positions, but is producing the palatal glide /j/ in place of the voiced palatal stop /ɟ/ in some instances. However, it appears ML is moving on to the first stage of acquisition as she is also producing the alveolar /t, d/ in place of palatal /c, ɟ/.

Stage One is represented by learners of Czech with some experience. Learners of Czech at this stage produce alveolar /t, d/ for both the alveolar stops /t, d/ and the palatal stops /c, ɟ/. They are able to perceptually contrast tokens of alveolar and palatal stops but are not yet contrasting them productively. Subject Two, JD appear to be at this stage: he scored higher than chance levels on the perception task but completely conflated the phonemic distinction between alveolar /t, d/ and palatal /c, ɟ/ in free production. In terms of Brown's model, learners of Czech at this stage appear to be accessing their L1 representations for coronal so that both /t, d/ and /c, ɟ/ are funneled into the alveolar category. At this stage, no new phonemic structure has been acquired. Flege's Speech Learning Model would cast this in terms of 'new' sounds versus 'similar' sounds, with palatal /c, ɟ/ being perceived as tokens of the similar sounding L1 alveolar stops /t, d/.

Native-speaker informants reported that when the English-speaking learner of Czech substitutes an alveolar stop in place of the palatal, meaning is generally retained. In terms of comprehension, they indicated that substitution of the alveolar for the palatal was greatly preferred over substitution of a palatal glide.

At Stage Two, learners appear to productively contrast alveolar /t, d/ with palatal /c, ɟ/ although the production is not yet fully native-like. Learners at Stage Two produce alveolar stop plus palatal glide sequences /tj, dj/ in place of palatal stops to maintain a phonemic contrast between alveolar /t, d/ and palatal /c, ɟ/. Native-speaker informants reported that /tj, dj/ sound substantially longer than a simple palatal segment [c, ɟ] and they classify it as a t + j sound.

Although none of the four subjects for whom I was able to gather production data were able to produce what could be considered a native-like palatal stop (by producing instead the sequences alveolar stop plus palatal glide), under the phonological model of L1 interference developed in this thesis, native English speakers should be able to eventually develop new segmental representations for the palatal stops /c, ɟ/ as they have the required featural "building block" of [posterior] with which to do so. I propose, then, an eventual third stage of acquisition for native English speakers:

Stage Three: Czech palatal stop /c, ɟ/

Able to productively contrast alveolar /t, d/ versus palatal /c, ɟ/

Unfortunately, none of the four subjects tested in the production experiments reached this third level of acquisition. However, if we accept the arguments presented in this thesis that (i), the ability to establish new segmental representations for a novel L2 segmental contrast relies on the L1 inventory of features, and (ii), English speakers do have the underlying feature [posterior] with which to contrast the novel segmental stop contrast of alveolar /t, d/ versus palatal /c, ɟ/, then we are faced with two questions. First, why were the speakers in the study unable to produce a native sounding Czech palatal stop? Can they be taught to produce true palatal stops [c, ɟ] and not an alveolar stop plus glide sequence?

Although none of the four subjects I tested had acquired a native-sounding palatal stop pronunciation, I have encountered native English speakers who are able to pronounce palatal /c, ɟ/ rather than the alveolar plus glide sequence produced by my subjects. These speakers were highly motivated and used Czech as a language of communication. As I noted in chapter Five, the subjects I tested did not use Czech as their language of communication and did not appear highly motivated. The move from Stage Two, where alveolars and palatals are in contrast (albeit in a nonnative manner), to Stage Three, where alveolars and native sounding palatals are in contrast may thus be motivated by extra-linguistic factors such as motivation and identification with the target language group.

The second question is: Why were the subjects' productions of alveolar plus glide sequences inconsistent in that they produced both alveolar /t, d/ and alveolar plus glide sequences /tj, dj/ for the palatal stop /c, j/? I attribute this non-native production pattern to two factors: orthography and morpho-phonemic considerations. As discussed in Chapter Four, Czech uses the symbols 't' and 'd' to indicate both alveolar /t, d/ and palatal /c, j/ with the palatal being distinguished from the alveolar by means of three orthographic conventions, depending on the following vowel quality.

The acquisition of the Czech alveolar versus palatal distinction may be further complicated by morphophonemic patterns, which may be misleading. Consider for example the difference between nominative versus locative forms in words ending with an alveolar consonant shown in (19):

(19)	<i>Nominative</i>		<i>Locative</i>	
(a)	třída	[tɾi:da]	'classroom'	ve třídě [fɾi:jɛ] 'in the classroom'
(b)	sešit	[sɛʃit]	'notebook'	v sešitě [fseʃice] 'in the notebook'
(c)	obchod	[opxot]	'store'	v obchodě [vopxoʃɛ] 'in the store'

Because a lexical item may have an alveolar in its nominative form and a palatal in the locative, native speakers are likely to understand through context what nonnative Czech speaker is referring to even if he pronounces an alveolar stop in place of a palatal. If communication is not hampered by the nonnative speakers' misuse of the alveolar stop in place of the palatal, he is not likely to change his pronunciation.

In this chapter I presented the results of experiments designed to test native English speakers' perception and production of the Czech alveolar /t, d/ versus palatal /c, j/ contrast.

Results showed that English speakers can perceptually contrast Czech alveolar versus palatal stops at greater than chance levels, but that production varies as a result of experience and length of exposure to Czech. These results generally supported Brown's model of L1 interference in L2 phonological acquisition discussed extensively in this

thesis. However, while Brown's model makes broad predictions as to the acquisition of a novel phonemic contrast based on the presence or absence of a contrasting feature in the L1 feature inventory, I found that the learners of Czech I studied varied in their production depending on length of exposure. Based on these varied yet systematic productions, I proposed a series of three stages of acquisition of Czech palatal stops by native English speakers.

CHAPTER SIX

AVENUES FOR FUTURE RESEARCH AND CONCLUSION

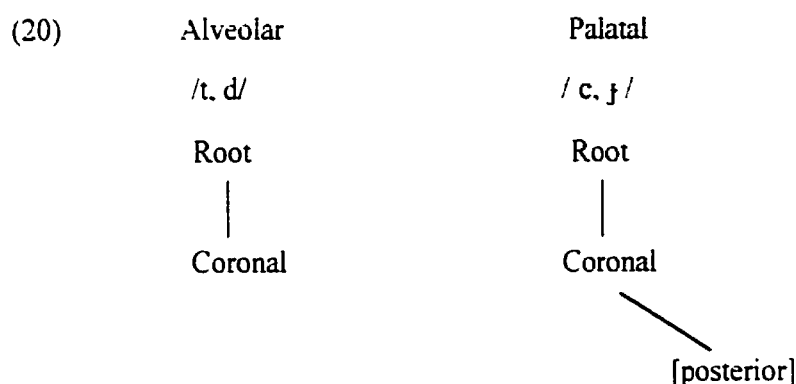
0.0 BACKGROUND

In this final chapter I will discuss avenues for future research and conclusions.

1.0 AVENUES FOR FUTURE RESEARCH

The results of the perception and production experiments reported in this thesis open up many avenues for future research. One such avenue is motivated by the limited pool of subjects tested: it would be illuminating to study fluent nonnative speakers of Czech as to their production of palatal stops. The subjects I tested here self-reported that, despite living in a Czech speaking environment, English was their main language of communication; they had little external motivation to learn Czech as their work and social communities were English speaking. I would predict that motivated speakers with more experience would be able to produce native-like Czech palatal stops.

Another avenue of research suggested by the assumptions and research in this thesis would be to test children learning Czech as their L1. In the model of segmental acquisition as a process of structure elaboration discussed in Chapter Two, Section 2.0, children learning language elaborate the UG-provided, minimal structure on a language-specific basis in a step-by-step manner following relationships of dependency and constituency encoded in the structure. The goal of the language learner is to expand the minimal language-specific structure until all the features that are necessary to distinguish all phonemes in the ambient language are present. Thus, we would predict that children learning Czech as their L1 should acquire the segmental structure for alveolar /t, d/ before the more structurally complex palatal /c, ɟ/. Recall the feature structures presented in Chapter Four for the alveolar versus palatal contrast, given here in (20):



While I was not able to experimentally test this assumption, I did have the opportunity to informally observe a nearly three-year-old child (35 months) over a period of a week. This child was producing alveolar /t, d/ segments in place of palatal stops /c, ʃ/, as would be predicted by the model of structural elaboration. Once contrast is detected in the input, he will elaborate the Coronal Place node to contrast alveolar /t, d/ with palatal /c, ʃ/. Further research on children's acquisition of Czech would be illuminating.

2.0 CONCLUSION

This thesis combined the fields of theoretical phonology and second language acquisition (specifically, second language phonology). Following Brown (1993: 1997), I presented the hypothesis that the phonological acquisition of novel L2 segmental contrasts is mediated by the system of features, rather than the segments per se, which are found in the L1.

As background, in Chapter One I presented the theories of Feature Geometry and Underspecification in support of the claim that the acquisition of segmental structure in both first and second language acquisition is a process of structure building.

Segmental structure building was introduced in Chapter Two on First Language Acquisition, where features are added on a step-by-step basis following relationships of constituency and dependency encoded in the hierarchy. Because features are added on a node-by-node basis in response to contrasts detected in the input, I argued that features

are monovalent, or privative.

As described in Chapter Three, the acquisition of the feature geometry of the first language is claimed to influence the acquisition of the L2 phonology. Thus, I argued for a Partial-Access approach to Universal Grammar, where UG is thought to be accessible to the L2 learner only in the form of parameters (here, the features) instantiated in the L1. The model developed in this thesis argues that if a particular feature is present in anywhere in the L1 feature hierarchy, then any novel L2 phonological contrast based on that feature can be acquired. Perceptual experience with the acoustic dimensions defined by a particular feature allow the learner to accurately discriminate any contrast based on that feature. However, if the L1 does not contrast a particular feature, perception of the novel L2 contrast will be precluded. Thus, UG is partially accessible to the L2 learner in the form of the features present in the L1. If UG was fully accessible to L2 learners, we would expect learners to be able to acquire *any* novel phonemic contrast. I presented work by Brown (1998) and Matthews (1997) showing that Japanese speakers have difficulty, even after training, in acquiring the English liquid contrast /l~r/ because they lack the requisite feature in their inventory.

In Chapter Four I argued that the contrasting feature for the Czech alveolar versus palatal distinction is the feature [posterior], and that this feature is present in English to contrast the alveolar fricatives /s, z/ with the alveo-palatal fricatives /ʃ, ʒ/.

In Chapter Five I presented the perception and production experiments I conducted to test the hypothesis that English speakers should be able to contrast Czech alveolar /t, d/ versus palatal /c, j/ since English has the featural building block of [posterior] in its inventory. Results showed that, (i), subjects could perceive the phonemic contrast between Czech alveolar /t, d/ and palatal /c, j/ at greater than chance levels; (ii), perceptual acuity increased with length of exposure to Czech; and (iii), subjects who were not absolute beginners were able to productively contrast alveolar /t, d/ and palatal /c, j/, although their productions of palatal stops were not fully native-like. Based on the subjects' varied yet systematic productions, I proposed a three-stage outline for the acquisition of the Czech palatal stops /c, j/ by native English speakers in which

subjects at the third stage are able to produce native-sounding palatal /c, j/. The likelihood of attaining this stage appeared to be influenced by a number of external factors such as learner motivation, reliance on the L1, and the learners' identification with the target language community.

In Chapter Five I also found that performance on the reading task was highly influenced by orthography. Czech has a system of diacritics indicating the presence of an alveolar or palatal segment. The results of the sentence reading task showed that the early learners of Czech did not appear to be sensitive to these orthographic cues and thus produced palatal segments as alveolars even when they were able to produce a (nonnative) sounding palatal in the same lexical item in free production. I thus argued that, although the sentence reading task was a test of orthographic knowledge and not of phonology, the unfamiliar Czech orthography can cause confusion for nonnative speakers and may negatively influence their contrastive production of the alveolar versus palatal segments.

APPENDIX ONE: FORCED CHOICE PHONEME SELECTION TASK
TEST STIMULI – ORTHOGRAPHIC FORM

1. duben	26. matěj	51. deka	76. dýmal
2. mladá	27. někdo	52. havet'	77. t'apka
3. kost	28. důkaz	53. t'oplá	78. hodil
4. t'opka	29. debil	54. doma	79. ud'ob
5. bašta	30. d'aha	55. d'ábel	80. plet'
6. tomu	31. tebe	56. kvílet	81. t'atka
7. schody	32. jekot	57. taky	82. jídám
8. doba	33. tyká	58. hut'	83. nosit
9. tolik	34. díra	59. tužka	84. beton
10. nejdu	35. děkan	60. mad'ar	85. zat'
11. fot'ák	36. t'ukám	61. potom	86. d'obal
12. t'íše	37. čedok	62. klonit	87. piti
13. prd'uch	38. tisíc	63. tukan	88. daleko
14. t'uhyl	39. pod'ub	64. sít'	89. pot'oh
15. prut	40. bělat	65. patý	90. smrad'och
16. d'ubám	41. jetel	66. rod'ák	91. vid'
17. tanec	42. stud	67. nat'uk	92. bodej
18. dělej	43. těžký	68. štuk	93. chata
19. teze	44. tykev	69. d'obu	94. podiv
20. rodu	45. nadej	70. jader	95. nastup
21. tělo	46. bat'ule	71. ojet	96. divák
22. hudě	47. kudy	72. věty	97. pěti
23. zírat	48. bat'oh	73. loukot'	98. dykal
24. sut'	49. chot'	74. d'ubal	99. ket'as
25. otěz	50. páteř	75. dárek	100. bat'

APPENDIX ONE
FORCED CHOICE PHONEME SELECTION TASK
TEST STIMULI – PHONETIC FORM

1. [duben]	26. [macej]	51. [deka]	76. [di:mal]
2. [mlada :]	27. [negdo]	52. [havjet]	77. [capka]
3. [kost]	28. [dukaz]	53. [copla:]	78. [hojl]
4. [copka]	29. [debil]	54. [doma]	79. [ujob]
5. [bafta]	30. [jaha]	55. [ja:bel]	80. [plec]
6. [tomu]	31. [tebe]	56. [kvi:let]	81. [cafka]
7. [sxodi]	32. [jekot]	57. [takı]	82. [jirda:m]
8. [doba]	33. [tıka:]	58. [huc]	83. [nosit]
9. [tolik]	34. [jira]	59. [tuзка]	84. [beton]
10. [nejdu]	35. [jekan]	60. [majar]	85. [zac]
11. [foca:k]	36. [cuka:m]	61. [potom]	86. [jobal]
12. [cise]	37. [tʃedok]	62. [klonit]	87. [pici:]
13. [prjux]	38. [cısr:ts]	63. [tukan]	88. [daleko]
14. [cuhi:l]	39. [pojub]	64. [si:c]	89. [pocox]
15. [prut]	40. [bjelat]	65. [pati:]	90. [smrajox]
16. [juba:m]	41. [jetel]	66. [rojak]	91. [vic]
17. [tanec]	42. [stut]	67. [nacuk]	92. [bodej]
18. [jelej]	43. [cezki:]	68. [stuk]	93. [xata]
19. [teze]	44. [tıkev]	69. [jobu]	94. [pojrv]
20. [rodu]	45. [najej]	70. [jader]	95. [nastup]
21. [celo]	46. [batule]	71. [ojet]	96. [jrva:k]
22. [huje]	47. [kudi]	72. [vjetı]	97. [peci:]
23. [zırat]	48. [bacox]	73. [loukoc]	98. [dıkal]
24. [suc]	49. [xoc]	74. [jubal]	99. [kecas]
25. [oce3]	50. [parter]	75. [darek]	100. [bac]

APPENDIX ONE: FORCED CHOICE PHONEME SELECTION TASK
TEST STIMULI: BREAKDOWN BY VOWEL AND WORD POSITION

WORD-INITIAL				
/ta/	/te/	/ti/	/to/	/tu/
17. tanec	19. teze	33. tyka	6. tomu	59. tužka
57. taky	31. tebe	44. tykev	9. tolik	63. tukan
/da/	/de/	/di/	/do/	/du/
75. dárek	29. debil	76. dýmal	8. doba	1. duben
88. daleko	51. deka	98. dykal	54. doma	28. důkaz
/ca/	/ce/	/ci/	/co/	/cu/
77. t'apka	21. tělo	12. t'íše	4. t'opka	14. t'uhyl
81. t'afka	43. těžky	38. tisíc	53. t'oplá	36. t'ukám
/ja/	/je/	/ji/	/jo/	/ju/
30. d'aha	18. dělej	34. díra	69. d'obu	16. d'ubám
55. d'abel	35. děkan	96. divák	86. d'obal	74. d'ubal
WORD-MEDIAL				
/ta/	/te/	/ti/	/to/	/tu/
5. bašta	41. jetel	65. patý	61. potom	68. štuk
93. chata	50. páteř	72. věty	84. beton	95. nastup
/da/	/de/	/di/	/do/	/du/
2. mladá	70. jader	7. schody	27. někdo	10. nejdu
82. jítam	92. bodej	47. kudy	37. čedok	20. rodu
/ca/	/ce/	/ci/	/co/	/cu/
11. foťák	25. otěž	87. pítí	48. bat'oh	46. bat'ule
99. ket'as	26. matěj	97. petí	89. pot'oh	67. nat'uk
/ja/	/je/	/ji/	/jo/	/ju/
60. maďar	22. huďe	78. hoďil	79. ud'ob	13. prďuch
66. roďák	45. naďej	94. poďiv	90. smrad'och	39. poďub
WORD-FINAL				
/at/	/et/	/it/	/ot/	/ut/
23. zírat	56. kvílet	62. klonit	3. kost	15. prut
40. bělat	71. ojet	83. nosit	32. jekot	42. stut
/ac/	/ec/	/ic/	/oc/	/uc/
85. zat'	52. ha'vet'	64. siť	49. choť	24. suť
100. bat'	80. plet'	91. víď	73. loukot'	58. huť

APPENDIX ONE
FORCED CHOICE PHONEME SELECTION TASK
TEST STIMULI

- | | | | |
|--------------|--------------|--------------|----------------|
| 1. __uben | 26. ma__ej | 51. __eka | 76. __ýmal |
| 2. mla__a | 27. nek__o | 52. have__ | 77. __apka |
| 3. kos__ | 28. __ukaz | 53. __oplá | 78. ho__il |
| 4. __opka | 29. __ebil | 54. __oma | 79. u__ob |
| 5. baš__a | 30. __aha | 55. __abel | 80. ple__ |
| 6. __omu | 31. __ebe | 56. kvile__ | 81. __afka |
| 7. scho__y | 32. jeko__ | 57. __aky | 82. ji__ám |
| 8. __oba | 33. __yká | 58. hu__ | 83. nosi__ |
| 9. __olik | 34. __íra | 59. __užka | 84. be__on |
| 10. nej__u | 35. __ekan | 60. ma__ar | 85. za__ |
| 11. fo__ak | 36. __ukam | 61. po__om | 86. __obal |
| 12. __íše | 37. če__ok | 62. kloni__ | 87. pi__i |
| 13. pr__uch | 38. __isic | 63. __ukan | 88. __aleko |
| 14. __uhyl | 39. po__ub | 64. si__ | 89. po__oh |
| 15. pru__ | 40. bela__ | 65. pa__ý | 90. smra__och |
| 16. __ubám | 41. je__el | 66. ro__ák | 91. vi__ |
| 17. __anec | 42. stu__ | 67. na__uk | 92. bo__ej |
| 18. __elaj | 43. __ezky | 68. š__uk | 93. cha__a |
| 19. __eze | 44. __ykev | 69. __obu | 94. po__iv |
| 20. ro__u | 45. na__ej | 70. ja__er | 95. nas__up |
| 21. __elo | 46. ba__ule | 71. oje__ | 96. __ivák |
| 22. hu__e | 47. ku__v | 72. ve__v | 97. pe__v |
| 23. zíra__ | 48. ba__oh | 73. louko__ | 98. __ykal |
| 24. su__ | 49. cho__ | 74. __ubal | 99. ke__as |
| 25. o__ež | 50. pa__eř | 75. __árek | 100. ba__ |

APPENDIX TWO
FORCED CHOICE PHONEME SELECTION TASK
RESULTS

SUBJECT ONE: ML

(** indicates incorrect response)

1. __uben	26. ma__ej	51. __eka	76. __ymal
2. mla__a	27. nek__o	**52. have__t_	77. __apka
3. kos__	28. __ukaz	53. __oplá	**78. ho_d_il
**4. _t_opka	29. __ebil	54. __oma	79. u__ob
5. baš__a	30. __aha	55. __abel	**80. ple__t_
6. __omu	31. __ebe	56. kvile__	**81. _t_afka
7. scho__y	32. jeko__	57. __aky	82. ji__ám
8. __oba	33. __vká	58. hu__	83. nosi__
9. __olik	**34. _d_íra	59. __užka	84. be__on
10. nej__u	35. __ekan	60. ma__ar	**85. za__t_
**11. fo__t_ak	36. __ukam	61. po__om	86. __obal
**12. _t__ise	37. če__ok	62. kloni__	87. pi__i
**13. pr__t_uch	**38. _t__isic	63. __ukan	88. __aleko
14. __uhyl	39. po__ub	**64. si__t_	89. po__oh
15. pru__	40. bela__	65. pa__y	90. smra__och
16. __ubam	41. je__el	66. ro__ak	**91. vi__t_
17. __anec	42. stu__	67. na__uk	92. bo__ej
18. __elaj	43. __ezky	68. š__uk	93. cha__a
19. __eze	44. __ykev	69. __obu	**94. po_d_iv
20. ro__u	45. na__ej	70. ja__er	95. nas__up
21. __elo	46. ba__ule	71. oje__	**96. _d__ivak
**22. hu_d__e	47. ku__y	72. ve__y	**97. pe__t_y
23. zira__	48. ba__oh	**73. louko__t_	98. __vkal
24. su__	**49. cho__t_	74. __ubal	99. ke__as
25. o__ež	50. pa__eř	75. __arek	**100. ba__t_

APPENDIX TWO

FORCED CHOICE PHONEME SELECTION TASK

SUBJECT TWO: JD

(** indicates incorrect response)

- | | | | |
|----------------------|----------------------|-----------------------|------------------------|
| 1. __uben | **26. ma_t_ej | 51. __eka | 76. __ymal |
| 2. mla__a | **27. nek_t_o | **52. have_t__ | **77. __t__apka |
| 3. kos__ | 28. __ukaz | 53. __oplá | 78. ho__il |
| 4. __opka | 29. __ebil | 54. __oma | 79. u__ob |
| 5. baš__a | 30. __aha | 55. __abel | **80. ple_t__ |
| 6. __omu | 31. __ebe | 56. kvile__ | 81. __afka |
| 7. scho__y | 32. jeko__ | 57. __aky | 82. ji__ám |
| 8. __oba | 33. __yká | **58. hu_t__ | 83. nosi__ |
| 9. __olik | 34. __íra | 59. __užka | 84. be__on |
| 10. nej__u | 35. __ekan | 60. ma__ar | **85. za_t__ |
| **11. fo_t_ak | 36. __ukam | 61. po__om | 86. __obal |
| 12. __ise | 37. če__ok | 62. kloni__ | **87. pi_t__i |
| 13. pr__uch | 38. __isic | 63. __ukan | 88. __aleko |
| 14. __uhyl | **39. po_d_ub | **64. si_t__ | 89. po__oh |
| 15. pru__ | 40. bela__ | 65. pa__y | 90. smra__och |
| 16. __ubam | 41. je__el | 66. ro__ak | 91. vi__ |
| 17. __anec | 42. stu__ | 67. na__uk | 92. bo__ej |
| 18. __elaj | 43. __ezky | 68. š__uk | 93. cha__a |
| 19. __eze | 44. __ykev | 69. __obu | **94. po_d_iv |
| 20. ro__u | 45. na__ej | 70. ja__er | 95. nas__up |
| 21. __elo | 46. ba__ule | 71. oje__ | **96. _d_ivak |
| 22. hu__e | **47. ku_t_y | 72. ve__y | 97. pe__y |
| 23. zira__ | 48. ba__oh | 73. louko__ | 98. __ykal |
| 24. su__ | **49. cho_t__ | 74. __ubal | **99. ke_t__as |
| 25. o__ež | 50. pa__eř | 75. __arek | **100. ba_t__ |

APPENDIX TWO

FORCED CHOICE PHONEME SELECTION TASK

SUBJECT THREE: AD

(** indicates incorrect response)

- | | | | |
|-------------------------|------------------------|-----------------------|------------------------|
| 1. __uben | 26. ma__ej | 51. __eka | 76. __ymal |
| 2. mla__a | 27. nek__o | **52. have__t_ | 77. __apka |
| 3. kos__ | 28. __ukaz | **53. _t__oplá | 78. ho__il |
| 4. __opka | 29. __ebil | 54. __oma | 79. u__ob |
| 5. baš__a | 30. __aha | 55. __abel | **80. ple__t__ |
| 6. __omu | 31. __ebe | 56. kvile__ | 81. __afka |
| 7. scho__y | 32. jeko__ | 57. __aky | 82. ji__ám |
| 8. __oba | 33. __yká | 58. hu__ | 83. nosi__ |
| 9. __olik | **34. _d__íra | 59. __užka | 84. be__on |
| 10. nej__u | 35. __ekan | 60. ma__ar | 85. za__ |
| 11. fo__ak | 36. __ukam | 61. po__om | 86. __obal |
| **12. _t__ise | 37. če__ok | 62. kloni__ | 87. pi__i |
| **13. pr__d__uch | 38. __isic | 63. __ukan | 88. __aleko |
| 14. __uhyl | **39. po__d__ub | **64. si__t__ | 89. po__oh |
| 15. pru__ | 40. bela__ | 65. pa__v | 90. smra__och |
| 16. __ubam | 41. je__el | 66. ro__ak | **91. vi__t__ |
| 17. __anec | 42. stu__ | 67. na__uk | 92. bo__ej |
| 18. __elaj | 43. __ezky | 68. š__uk | 93. cha__a |
| 19. __eze | 44. __ykev | 69. __obu | **94. po__d__iv |
| 20. ro__u | 45. na__ej | 70. ja__er | 95. nas__up |
| 21. __elo | 46. ba__ule | 71. oje__ | **96. _d__ivak |
| 22. hu__e | 47. ku__y | 72. ve__y | **97. pe__t__y |
| 23. zira__ | 48. ba__oh | 73. louko__ | 98. __ykal |
| 24. su__ | 49. cho__ | 74. __ubaí | 99. ke__as |
| 25. o__ež | 50. pa__eř | 75. __arek | **100. ba__t__ |

APPENDIX TWO

FORCED CHOICE PHONEME SELECTION TASK

SUBJECT FOUR: SW

(** indicates incorrect response)

- | | | | |
|----------------------|-----------------------|---------------|-----------------------|
| 1. __uben | 26. ma__ej | 51. __eka | 76. __ymal |
| 2. mla__a | 27. nek__o | 52. have__ | 77. __apka |
| 3. kos__ | 28. __ukaz | 53. __oplá | 78. ho__il |
| 4. __opka | 29. __ebil | 54. __oma | **79. u__d__ob |
| 5. baš__a | 30. __aha | 55. __abel | **80. ple__t__ |
| 6. __omu | 31. __ebe | 56. kvile__ | 81. __afka |
| 7. scho__y | 32. jeko__ | 57. __aky | 82. ji__ám |
| 8. __oba | 33. __yká | 58. hu__ | 83. nosi__ |
| 9. __olik | 34. __íra | 59. __užka | 84. be__on |
| 10. nej__u | 35. __ekan | 60. ma__ar | 85. za__ |
| 11. fö__ak | **36. _t__ukam | 61. po__om | 86. __obal |
| 12. __ise | 37. če__ok | 62. kloni__ | **87. pi__t__i |
| 13. pr__uch | **38. _t__isic | 63. __ukan | 88. __aleko |
| 14. __uhyl | 39. po__ub | 64. si__ | 89. po__oh |
| 15. pru__ | 40. bela__ | 65. pa__y | 90. smra__och |
| 16. __ubam | 41. je__el | 66. ro__ak | **91. vi__t__ |
| 17. __anec | 42. stu__ | 67. na__uk | 92. bo__ej |
| 18. __elaj | 43. __ezky | 68. š__uk | 93. cha__a |
| 19. __eze | 44. __ykev | 69. __obu | 94. po__iv |
| 20. ro__u | 45. na__ej | 70. ja__er | 95. nas__up |
| 21. __elo | 46. ba__ule | 71. oje__ | **96. _d__ivak |
| 22. hu__e | 47. ku__y | 72. ve__y | 97. pe__y |
| 23. zira__ | 48. ba__oh | 73. louko__ | 98. __ykal |
| **24. su__t__ | 49. cho__ | 74. __ubal | 99. ke__as |
| 25. o__ež | 50. pa__eř | 75. __arek | 100. ba__ |

APPENDIX TWO

FORCED CHOICE PHONEME SELECTION TASK

SUBJECT FIVE: JA

(** indicates incorrect response)

- | | | | |
|-----------------------|----------------------|-----------------------|-------------------------|
| 1. __uben | 26. ma__ej | 51. __eka | 76. __ymal |
| 2. mla__a | 27. nek__o | **52. have__t_ | 77. __apka |
| 3. kos__ | 28. __ukaz | 53. __oplá | 78. ho__il |
| 4. __opka | 29. __ebil | 54. __oma | **79. u__d_ob |
| 5. baš__a | 30. __aha | 55. __abel | **80. ple__t_ |
| 6. __omu | 31. __ebe | 56. kvile__ | 81. __afka |
| 7. scho__v | 32. jeko__ | 57. __aky | 82. ji__ám |
| 8. __oba | 33. __yká | 58. hu__ | 83. nosi__ |
| 9. __olik | **34. __d_íra | 59. __užka | 84. be__on |
| 10. nej__u | 35. __ekan | 60. ma__ar | 85. za__ |
| 11. fo__ak | **36. _t_ukam | 61. po__om | 86. __obal |
| **12. _t_ise | 37. če__ok | 62. kloni__ | 87. pi__i |
| **13. pr_d_uch | 38. __isic | 63. __ukan | 88. __aleko |
| 14. __uhyl | 39. po__ub | 64. si__ | 89. po__oh |
| 15. pru__ | 40. bela__ | 65. pa__v | **90. smra_d_och |
| 16. __ubam | 41. je__ei | 66. ro__ak | **91. vi__t_ |
| 17. __anec | 42. stu__ | 67. na__uk | 92. bo__ej |
| 18. __elaj | 43. __ezky | 68. š__uk | 93. cha__a |
| 19. __eze | 44. __ykev | 69. __obu | 94. po__iv |
| 20. ro__u | 45. na__ej | 70. ja__er | 95. nas__up |
| 21. __elo | 46. ba__ule | 71. oje__ | **96. __t_ivak |
| 22. hu__e | 47. ku__v | 72. ve__v | 97. pe__v |
| 23. zira__ | 48. ba__oh | 73. louko__ | 98. __ykal |
| 24. su__ | 49. cho__ | 74. __ubal | 99. ke__as |
| 25. o__ež | 50. pa__eř | 75. __arek | 100. ba__ |

APPENDIX TWO
FORCED CHOICE PHONEME SELECTION TASK
RESULTS

SUBJECT SIX: RK

(** indicates incorrect response)

- | | | | |
|---------------------------|----------------------|----------------------|------------------------|
| 1. __uben | 26. ma__ej | 51. __eka | 76. __ymal |
| 2. mla__a | 27. nek__o | **52. have__t | 77. __apka |
| 3. kos__ | 28. __ukaz | 53. __oplá | 78. ho__il |
| 4. __opka | 29. __ebil | 54. __oma | 79. u__ob |
| 5. baš__a | 30. __aha | 55. __abel | **80. ple__t |
| 6. __omu | 31. __ebe | 56. kvile__ | 81. __afka |
| 7. scho__y | 32. jeko__ | 57. __aky | 82. ji__ám |
| 8. __oba | 33. __yká | 58. hu__ | 83. nosi__ |
| 9. __olik | **34. _d__íra | 59. __užka | 84. be__on |
| 10. nej__u | 35. __ekan | 60. ma__ar | 85. za__ |
| 11. fo__ak | 36. __ukam | 61. po__om | 86. __obal |
| 12. __ise | 37. če__ok | 62. kloni__ | 87. pi__i |
| 13. **pr__d__uch | 38. __isic | 63. __ukan | 88. __aleko |
| 14. __uhyl | 39. po__ub | 64. si__ | 89. po__oh |
| 15. pru__ | 40. bela__ | 65. pa__y | 90. smra__och |
| 16. __ubam | 41. je__el | 66. ro__ak | 91. vi__ |
| 17. __anec | 42. stu__ | 67. na__uk | 92. bo__ej |
| 18. __elaj | 43. __ezky | 68. š__uk | 93. cha__a |
| 19. __eze | 44. __ykev | 69. __obu | **94. po__d__iv |
| 20. ro__u | 45. na__ej | 70. ja__er | 95. nas__up |
| 21. __elo | 46. ba__ule | 71. oje__ | 96. __ivak |
| 22. hu__e | 47. ku__y | 72. ve__v | 97. pe__y |
| 23. zira__ | 48. ba__oh | 73. louko__ | 98. __ykal |
| 24. su__ | 49. cho__ | 74. __ubal | 99. ke__as |
| 25. o__ež | 50. pa__eř | 75. __arek | 100. ba__ |

APPENDIX THREE
EXPERIMENT TWO: PRODUCTION

STIMULI AND RESULTS
SENTENCE READING TASK

1. Tati. udělej to ted’.

[t|a|c|i. u|ʃ|ělej [t|o [t|e|c].

ML	[t]	[t]	[d]	[t]	[t]	[d]
JD	[t]	[t]	[d]	[t]	[t]	[d]
JA	[t]	[t]	[dj]	[t]	[t]	[d]
RK	[t]	[t]	[dj]	[t]	[t]	[t]

2. Toto telátko je hrdina rodinu.

[t|o[t|o [t|elá[t|ko je hr|ʃ|ina ro|ʃ|inu.

ML	[t]	[t]	[t]	[t]	[d]	[d]
JD	[t]	[t]	[t]	[t]	[d]	[d]
JA	[t]	[t]	[t]	[t]	[d]	[d]
RK	[t]	[t]	[t]	[t]	[d]	[d]

3. Viděla d’abel a nadávala ho.

Vi|ʃ|ěla |ʃ|abel a na|d|avala ho.

ML	[d]	[d]	[d]
JD	[d]	[d]	[d]
JA	[dj]	[d]	[d]
RK	[dj]	[d]	[d]

4. Co měla dělat, poděkovat ti?

Co měla [j]ěla[t], po[j]ěkova[t] [c]i?

ML	[d]	[t]	[d]	[t]	[t]
JD	[d]	[t]	[d]	[t]	[t]
JA	[dj]	[t]	[dj]	[t]	[t]
RK	[d]	[t]	[dj]	[t]	[t]

5. Ještě ti to nedošlo?

Ješ[c]ě [c]i [t]o ne[d]ošlo?

ML	[t]	[t]	[t]	[d]
JD	[t]	[t]	[t]	[d]
JA	[tj]	[t]	[t]	[d]
RK	[tj]	[t]	[t]	[d]

6. Dělam to díky tobě.

[j]ělam [t]o [d]íky [t]obě.

ML	[d]	[t]	[d]	[t]
JD	[d]	[t]	[d]	[t]
JA	[dj]	[t]	[d]	[t]
RK	[dj]	[t]	[d]	[t]

7. Věděla jsem, že přijdeš pozdě.

Vě[j]e&la jsem, že přij[d]eš pož[d]ě.

ML	[d]	[d]	[d]
JD	[d]	[d]	[d]
JA	[dj]	[d]	[dj]
RK	[dj]	[d]	[dj]

8. Vidím choť a rodinu.

Vi[ɟ]ím cho[c] a ro[ɟ]inu.

ML	[d]	[t]	[d]
JD	[d]	[t]	[d]
JA	[d]	[t]	[d]
RK	[d]	[t]	[d]

9. Vratím se ještě v středu.

Vra[c]ím se jes[c]ě v s[t]ře[d]u.

ML	[t]	[t]	[t]	[d]
JD	[t]	[t]	[t]	[d]
JA	[t]	[tj]	[t]	[d]
RK	[t]	[tj]	[t]	[d]

10. Otec je křestan, matka ne.

O[t]ec je křes[c]an, ma[t]ka ne.

ML	[t]	[t]	[t]
JD	[t]	[t]	[t]
JA	[t]	[t]	[t]
RK	[t]	[t]	[t]

11. Těším se davno na ticho.

[c]ěším se [d]avno na [c]icho.

ML	[t]	[d]	[t]
JD	[t]	[d]	[t]
JA	[tj]	[d]	[t]
RK	[t]	[d]	[t]

12. Děvče, dám ti vědět zítřa.

[ɟ]ěvče. [d]ám [c]i vě[ɟ]ě[t] zít[ɟ]ra.

ML	[d]	[d]	[t]	[d]	[t]	[t]
JD	[d]	[d]	[t]	[d]	[t]	[t]
JA	[dj]	[d]	[t]	[dj]	[t]	[t]
RK	[dj]	[d]	[t]	[dj]	[t]	[t]

13. Tyden budete chovat děti.

[t]y[d]en bu[d]e[t]e chova[t] [ɟ]ě[c]i.

ML	[t]	[d]	[d]	[t]	[t]	[d]	[t]
JD	[t]	[d]	[d]	[t]	[t]	[d]	[t]
JA	[t]	[d]	[d]	[t]	[t]	[dj]	[t]
RK	[t]	[d]	[d]	[t]	[t]	[dj]	[t]

14. Sedíme a díváme v divadle.

Se[ɟ]íme a [ɟ]iváme v [ɟ]iva[d]le.

ML	[d]	[d]	[d]	[d]
JD	[d]	[d]	[d]	[d]
JA	[d]	[d]	[d]	[d]
RK	[d]	[d]	[dj]	[d]

15. Teď je tři čtvrtě na deset.

[t]e[c] je [t]ři č[t]vr[c]ě na [d]ese[t].

ML	[t]	[d]	[t]	[t]	[t]	[d]	[t]
JD	[t]	[d]	[t]	[t]	[t]	[d]	[t]
JA	[t]	[d]	[t]	[t]	[tj]	[d]	[t]
RK	[t]	[d]	[t]	[t]	[t]	[d]	[t]

REFERENCES

- Abramson, A., & Lisker, K. (1970). Discriminability along the voicing continuum: Cross-language tests. In B. Hála, M. Romportal, & P. Janota. (Eds.). *Proceedings of the 6th International Congress of Phonetic Sciences*. (pp. 569-573). Prague: Academia.
- Anderson, J., & Ewen, C. (1987). *Principles of dependency phonology*. Cambridge: Cambridge University Press.
- Archangeli, D. (1988). Aspects of underspecification theory. *Phonology*, 5, 183-207.
- Archibald, J.A. (1992). Transfer of L1 parameter settings: Some empirical evidence from Polish metrics. *Canadian Journal of Linguistics*, 37, 301-339.
- Archibald, J.A. (1993). *Language learnability and L2 phonology: The acquisition of metrical parameters*. Dordrecht: Kluwer.
- Arlt, P., & Goodman, M. (1976). A comparative study of articulation acquisition as based on a study of 240 normals, aged three to six. *Language, Speech and Hearing Services in Schools*, 7, 173-180.
- Atkinson, M. (1982). *Explanations in the study of child language development*. Cambridge: Cambridge University Press.
- Avery, P., & Rice, K. (1989). Segment structure and coronal underspecification. *Phonology*, 6 (2), 179-200.
- Berwick, R.C. (1985). *The acquisition of syntactic knowledge*. Cambridge, MA: MIT Press.
- Best, P. (1993). Emergence of language-specific constraints in perception of non-native speech: A window on early phonological development. In B. de Boysson-Bardies. (Ed.). *Developmental Neurocognition: Speech and Face Processing in the First Year of Life*. (pp. 289-304). Dordrecht: Kluwer Academic Publishers.
- Best, P. (1994). The emergence of native-language phonological influence on infants: A perceptual assimilation model. In H. Nusbaum, J. Goodman, & C. Howard. (Eds.). *The Transition from Speech Sounds to Spoken Words: The Development of Speech Perception*. (pp. 167-224). Cambridge, MA: MIT Press.
- Bley-Vroman, R. (1990). The logical problem of second language learning. *Linguistic Analysis*, 20, 3-49.

- Bluhme, H. (1969). Zur Perzeption deutscher Sprachlaute durch englischsprechende Studenten. *Phonetica*, 20, 57-62.
- Brennan, E.M., & Brennan, J.S. (1981). Accent scaling and language attitudes: Reactions to Mexican-American English speech. *Language and Speech*, 24, 207-222.
- Broselow, E., & Finer, D. (1991). Parameter setting in second language phonology and syntax. *Second Language Research*, 7, 35-59.
- Brown, C. (1993). The role of the L1 grammar in the L2 acquisition of segmental structure. *McGill Working Papers in Linguistics*, 9, 180-210.
- Brown, C. (1997). Acquisition of segmental structure. PhD Thesis. McGill University.
- Brown, C. (1998). The role of the L1 grammar in the L2 acquisition of segmental structure. *Second Language Research*, 14 (2), 136-193.
- Brown, C. (2000). The interrelationship between the L1 feature inventory and the acquisition of L2 segmental contrasts. In J. Archibald, (Ed.), *Second Language Acquisition and Linguistic Theory*. Blackwell.
- Brown, C., & Matthews, J. (1993). The acquisition of segmental structure. *McGill Working Papers in Linguistics*, 9, 180-210.
- Brown, C., & Matthews, J. (1997). The role of feature geometry in the development of phonemic contrasts. In S.J. Hannah & M. Young-Scholten, (Eds.), *Focus on Phonological Acquisition* Philadelphia: John Benjamins.
- Brown, R., & Hanlon, C. (1970). Derivational complexity and the order of acquisition in child speech. In J.R. Hayes, (Ed.), *Cognition and the Development of Language*. (pp. 155-207). New York: Wiley.
- Carlton, T. (1991). *Introduction to the phonological history of the Slavic languages*. Columbus, Ohio: Slavica.
- Carroll, S. (1999). Putting 'input' in its proper place. *Second Language Research*, 15 (4), 337-388.
- Chlumsky, J. (1941). La photographie des articulations dessinées au palais artificiel. *Revue de Phonétique*, 4, 46-58.
- Chomsky, N. (1975). *Reflections on language*. New York: Pantheon.

- Chomsky, N. (1981). Principles and parameters in syntactic theory. In N. Hornstein & D. Lightfoot. (Eds.), *Explanation in Linguistics: The Logical Problem of Language Acquisition*. London: Longman.
- Chomsky, N. (1988). *Language and problems: The Managua lectures*. Cambridge: MIT Press.
- Chomsky, N., & Halle, M. (1968). *The sound pattern of English*. New York: Harper & Row.
- Clahsen, H., Eisenbeiss, S. & Vainikka, A. (1994). The seeds of structure: A syntactic analysis of the acquisition of case marking. In T. Hoekstra & B.D. Schwartz. (Eds.), *Language Acquisition Studies in Generative Grammar*. (pp. 85-118). Amsterdam: John Benjamins.
- Clements, G.N. (1985). The geometry of phonological features. *Phonology*. 2. 225-252.
- Clements, G.N. (1988). Towards a substantive theory of feature specifications. *Proceedings of NELS*. 18. 79-93.
- Clements, G.N., & Hume, E. (1995). The internal organization of speech sounds. In J. Goldsmith. (Ed.), *The Handbook of Phonological Theory*. (pp. 245-306). Oxford: Basil Blackwell.
- Cook, V.J. (1988). *Chomsky's Universal Grammar*. Oxford: Basil Blackwell.
- Corder, S.P. (1967). The significance of learners' errors. *International Review of Applied Linguistics*. 5. 161-170.
- Dankovičová, J. (1999). Czech. In *The Handbook of the IPA*. (pp. 70-73). Cambridge: Cambridge University Press.
- Dinnsen, D. (1996). Context-sensitive underspecification and the acquisition of phonemic contrasts. *Journal of Child Language*. 23. 57-79.
- Dresher, B.E. (1981). On the learnability of abstract phonology. In C. Baker & J. McCarthy. (Eds.), *The Logical Problem of Language Acquisition*. (pp. 188-210). Cambridge, MA: MIT Press.
- Dresher, B.E. (1999). Charting the learning path: Cues to parameter setting. *Linguistic Inquiry*, 30 (1), 27-67.

- Dresher, B.E. & Kaye, J.D. (1990). A computational learning model for metrical phonology. *Cognition*, 34 (2), 137-195.
- Fee, E.J. (1995). Segments and syllables in early language acquisition. In J. Archibald, (Ed.), *Phonological Acquisition and Phonological Theory*. (pp. 23-42). Hillsdale: Erlbaum.
- Finer, D., & Broselow, E. (1986). Second language acquisition of reflexive-binding. In *Proceedings of NELS 16*, (pp. 154-168). Amherst: University of Massachusetts. Graduate Linguistics Students Association.
- Flege, J.E. (1987). Effects of equivalence classification on the production of foreign language speech sounds. In A. James & J. Leather, (Eds.), *Sound Patterns in Second Language Acquisition*, (pp. 9-39). Dordrecht: Foris.
- Flege, J.E. (1988). The production and perception of foreign language speech sounds. In H. Winitz, (Ed.), *Human Communication and its Disorders*. (pp. 244-401). Norwood: Ablex.
- Flege, J.E. (1990). English vowel production by Dutch talkers: More evidence for the "similar" vs. "new" distinction. In J. Leather & A. James, (Eds.), *New Sounds 90: Proceedings of the 1990 Amsterdam Symposium on the Acquisition of Second Language Speech*. (pp. 255-293). Amsterdam: University of Amsterdam.
- Flege, J.E. (1991). Perception and production: The relevance of phonetic input to L2 phonological learning. In T. Huebner & C. Ferguson, (Eds.), *Crosscurrents in Second Language Acquisition and Linguistic Theory*. (pp. 249-290). Philadelphia: John Benjamins Publishing.
- Flege, J.E. (1995). Second language speech learning: Theory, findings and problems. In W. Strange, (Ed.), *Speech Perception and Linguistic Experience: Issues in Cross-language Research*. (pp. 233-273). Baltimore: York Press.
- Flynn, S. (1987). Contrast and construction in a parameter-setting model of L2 acquisition. *Language Learning: A Journal of Applied Linguistics*, 37 (1), 19-62.
- Flynn, S. (1991). Government-binding: Parameter setting in second language acquisition. In C. Ferguson & T. Huebner, (Eds.), *Crosscurrents in Second Language Acquisition and Linguistic Theories*, (pp. 143-167). Amsterdam: John Benjamins.
- Flynn, S. (1993). Interactions between L2 acquisition and linguistic theory. In F. Eckman, (Ed.), *Confluence: Linguistics, L2 Acquisition and Speech Pathology*. (pp. 15-36). Amsterdam : Benjamins.

- Flynn, S., & Martohardjono, G. (1994). Mapping from the initial state to the final state: The separation of universal principles and language specific properties. In B. Lust, J. Whitman & J. Kornfilt, (Eds.), *Syntactic Theory and First Language Acquisition: Cross-linguistic Perspectives: Vol. 1 Phrase Structure*. Hillsdale: Erlbaum.
- Fodor, J. (1999). Learnability theory: Triggers for parsing with. In E. Klein & G. Martohardjono, (Eds.), *The Development of Second Language Grammars: A Generative Perspective*. John Benjamins.
- Goto, H. (1971). Auditory perception by normal Japanese adults of the sounds "l" and "r". *Neuropsychologia*, 9, 317-323.
- Gregg, K. (1996). The logical and developmental problems of second language acquisition. In W. Ritchie & T. Bhatia, (Eds.), *Handbook of Second Language Acquisition* (pp. 50-84). San Diego: Academic Press.
- Grunwell, P. (1982). *Clinical phonology*. London: Croom Helm.
- Gussenhoven, C., & Jacobs, H. (1998). *Understanding phonology*. New York: Oxford University Press.
- Hála, B. (1923). *K popisu pražské vyslovnosti* 56 (3). Prague: České Akademie Věd a Umění.
- Hála, B. (1962). *Uvedení do foneticky češtiny na obecně fonetickém základě*. Prague: Nakladatelství Československé Akademie Věd.
- Hall, T.A. (1997). *The Phonology of Coronals: Current Issues in Linguistic Theory* 149 Philadelphia: John Benjamins North America.
- Hancin-Bhatt, B. (1994). Segment transfer: A consequence of a dynamic system. *Second Language Research*, 10 (3), 241-269.
- Heyde, A. (1979). The relationship between self-esteem and the oral production of a second language. Dissertation Abstracts International, Ann Arbor, MI.
- Hornstein, N., & Lightfoot, D. (1981). (Eds.). *Explanation in linguistics: The logical problem of language acquisition*. London: Longman.
- Hume, E. (1992). Front vowels, coronal consonants and their interaction in non-linear Phonology. Ph.D. dissertation. Cornell University.
- Hume, E. (1996). Coronal consonant, front vowel parallels in Maltese. *Natural Language and Linguistic Theory*, 14, 163-203.

- Ingram, D. (1995). The acquisition of negative constraints, the OCP, and underspecified representations. In J. Archibald, (Ed.), *Phonological Acquisition and Phonological Theory*, (pp. 63-80). Hillsdale: Erlbaum.
- Inkelas, S. (1994). *The consequences of optimization for underspecification*. ms. University of California, Berkeley.
- Jakobson, R. (1968). *Child language, aphasia, and phonological universals*. The Hague: Mouton. (Original work published in 1941.)
- Jakobson, R., Fant, G. & Halle, M. (1952). *Preliminaries to Speech Analysis*. Cambridge: MIT Press.
- Kean, M.L. (1975). The theory of markedness in generative grammar. Doctoral Dissertation. MIT, Cambridge, Massachusetts. Distributed 1980. Indiana University Linguistics club, Bloomington.
- Keating, P.A. (1988). Palatals as complex segments: X-ray evidence. *UCLA Working Papers in Phonetics*, 69, 77-91
- Keating, P. A. (1991). In C. Paradis & J.F. Prunet, (Eds.), *The Special Status of Coronals: Internal and External Evidence. II: Phonetics and Phonology*, (pp. 29-48). San Diego: Academic.
- Keating, P.A., & Lahiri, A. (1993). Fronted velars, palatalized velars, and palatals. *Phonetica*, 50, 73-101.
- Kiparsky, P. (1982). Lexical morphology & phonology. In I.S Yang, (Ed.), *Linguistics in the Morning Calm*, (pp. 3-91). Seoul, Korea: Hanshin.
- Kiparsky, P. (1985). Some consequences of lexical phonology. *Phonology Yearbook*, 2, 85-138.
- Kučera, H. (1961). *The phonology of Czech*. The Hague: Mouton.
- Ladefoged, P. (1982). *A course in phonetics*. New York: Harcourt Brace Jovanovich.
- Ladefoged, P., & Maddieson, I. (1988). Phonological features for places of articulation. In L. Hyman & C. Li, (Eds.), *Language, Speech and Mind: Studies in Honour of Victoria Fromkin*. London: Routledge.
- Lahiri, A., & Blumstein, S.E. (1984). A re-evaluation of the feature coronal. *Journal of Phonetics*, 12 (2), 133-145.

- Leather, J., & James, A. (1996). Second language speech. In W. Ritchie & T. Bhatia, (Eds.), *Handbook of Second Language Acquisition*, (pp. 269-316). San Diego: Academic Press, Inc.
- Lightfoot, D. 1981. In N. Hornstein & D. Lightfoot, (Eds.), *Explanation in Linguistics: The Logical Problem of Language Acquisition*. London: Longman
- Locke, J.L. (1968). Oral perception and articulation learning. *Perceptual and Motor Skills*, 26, 1259-1264.
- Locke, J.L. (1969). Short-term memory, oral perception and experimental sound learning. *Journal of Speech and Hearing Research*, 12, 185-192.
- Long, M.H. (1993) . Assessment strategies for second language acquisition theories. *Applied Linguistics*, 14, 225-249.
- Lust, B. (1994). Functional projection of CP and phrase structure parametrization: An argument for the strong continuity hypothesis. In B. Lust, J. Whitman & M. Suner, (Eds.), *Crosslinguistic Perspectives. Volume 1: Heads, Projections and Learnability*. (pp. 85-118). Hillsdale, N.J.: Erlbaum.
- Maddieson, I. (1984). *Patterns of sounds*. Cambridge: Cambridge University Press.
- Maddieson, I. (1987). The Margi vowel system and labiodororals. *Studies in African-Linguistics*, 18 (3), 327-355.
- Manzini, M., & Wexler, K. (1987). Parameters, binding theory, and learnability. *Linguistic Inquiry*, 18 (3), 413-444.
- Martohardjono, G. (1991). Universal aspects in the acquisition of a second language. Paper presented at the 19th Boston University Conference on Language Development, Boston University.
- Martohardjono, G. (1993). Wh-movement in the acquisition of a second language: A cross-linguistic study of three languages with and without movement. Doctoral Dissertation, Cornell University, Ithaca, N.Y.
- Matthews, J. (1997). Influence of pronunciation training on the perception of second language contrasts. In J. Leather & A. James, (Eds.), *New Sounds 97: Proceedings of the 1990 Amsterdam Symposium on the Acquisition of Second Language Speech*. Amsterdam: University of Amsterdam.
- McCarthy, J. (1988). Feature geometry and dependency: A review. *Phonetica*, 45, 85-108.

- Odden, D. (1978). Further evidence for the feature [grave]. *Linguistic Inquiry*, 9, 141-144.
- Pagliuca, W., & Mowrey, R. (1980). On certain evidence for the feature [grave]. *Language*, 56, 503-14.
- Palková, Z. (1997). *Fonetika a fonologie Češtiny*. Prague: Karolinum.
- Paradis, C., & Prunet, J.F. (1989). On coronal transparency. *Phonology*, 6 (2), 317-348.
- Paradis, C., & Prunet, J.F. (1990). On explaining some OCP violations. *Linguistic Inquiry*, 21(3), 456-466
- Paradis, C., & Prunet, J.F., (Eds.). (1991). *The Special Status of Coronals: Internal and External Evidence, II: Phonetics and Phonology*. San Diego: Academic Press.
- Paradis, J., & Genesee, F. (1997). On continuity and the emergence of functional categories in bilingual first language acquisition. *Language Acquisition*, 6, 91-124.
- Pater, J. (1993). Theory and methodology in the study of metrical parameter (re)setting. *McGill Working Papers in Linguistics*, 9 (1-2), 211-43.
- Piggot, G. (1988). The parameters of nasalization. *McGill Working Papers in Linguistics*, 5 (2), 128-177.
- Pisoni, D. (1973). Auditory and phonetic memory codes in the discrimination of consonants and vowels. *Perception and Psychophysics*, 13, 253-250.
- Prather, E., Hedrick, D., & Kern, C. (1975). Articulation development in children aged two to four years. *Journal of Speech and Hearing Research*, 40, 179-191.
- Prince, A., & Smolensky, P. (1993). *Optimality theory: Constraint interaction in generative grammar*. ms. Rutgers University and University of Colorado.
- Purcell, E.T., & Suter, R.W. (1980). Predictors of pronunciation accuracy: A re-examination. *Language Learning*, 30, 271-288.
- Pye, C., Ingram, D., & List, H. (1987). A comparison of initial consonant acquisition in English and Quiché. In K. Nelson & A. Van Kleeck (Eds.), *Children's Language*, Vol 6, (pp. 175-190). Hillsdale: Lawrence Erlbaum.
- Recasens, D. (1990). The articulatory characteristics of palatal consonants. *Journal of Phonetics*, 18, 267-280.

- Repp, B. (1984). Categorical perception: Issues, methods, findings. In N.J. Lass, (Ed.), *Speech and Language: Advances in Basic Research and Practice, Vol. 10*. New York: Academic Press.
- Rice, K. (1989). Autosegmental processes. *Canadian Journal of Linguistics*, 34 (1), 59-76.
- Rice, K. (1996). Default variability: The coronal-velar relationship. *Natural Language and Linguistics*, 14, 493-543.
- Rice, K., & Avery, P. (1995). Variability in a deterministic model of language acquisition: A theory of segmental elaboration. In J. Archibald, (Ed.), *Phonological Acquisition and Phonological Theory*, (pp. 23-42). Hillsdale: Erlbaum.
- Roca, I. (1994). *Generative phonology*. London: Routledge.
- Sagey, B. (1986). The representations of features and relations in non-linear phonology. Doctoral Dissertation, MIT, Cambridge, Massachussets.
- Schacter, J. (1989). Testing a proposed universal. In S. Gass, & J. Schacter, (Eds.), *Linguistic Perspectives on Second Language Acquisition*, (pp. 73-88). Cambridge: Cambridge University Press.
- Schacter, J. (1996). Maturation and the issue of Universal Grammar in second language acquisition. In W. Ritchie & T. Bhatia, (Eds.), *Handbook of Second Language Acquisition* (pp. 159-194). San Diego: Academic Press.
- Schouten, (1975). Native language interference in the perception of second language vowels. Dissertation Abstracts International, Ann Arbour, MI, 37, 98C.
- Selinker, L. (1972). Interlanguage. *International Review of Applied Linguistics*, 10, 209-231.
- Sheldon, A., & Strange, W. (1982). The acquisition of /r/ and /l/ by Japanese learners of English: Evidence that speech production can precede speech perception. *Applied Psycholinguistics*, 3 (3), 243-261.
- Short, D. (1993) Czech. In B. Comrie & G. Corbett, (Eds.), *The Slavonic Languages*, (pp. 455-532). London: Routledge.
- Stemberger J.P., & Stoel-Gammon, C. (1989). Consonant harmony and underspecification in child phonology, ms. University of Minneapolis. Minneapolis.

- Steriade, D. (1987). Redundant values. *Papers from the Chicago Linguistics Society* 23. *Part II: Parasession on autosegmental and metrical phonology*. (pp. 339-362).
- Steriade, D. (1995). Underspecification and markedness. In J. Goldsmith. (Ed.). *The Handbook of Phonological Theory*. Cambridge: Blackwell.
- Stevens, K. (1972). Sources of inter- and intra-speaker variability in the acoustic properties of speech sounds. In A. Rigault et al., (Eds.), *Proceedings of the Seventh International Congress of Phonetic Sciences*. (pp. 206-32). The Hague: Mouton.
- Stevens, K. (1989). On the quantal nature of speech. *Journal of Phonetics*, 17 (1-2), 3-45.
- Stoel-Gammon, C. (1985). Phonetic inventories, 15-24: A longitudinal study. *Journal of Speech and Hearing Research*, 28, 505-512.
- Stoel-Gammon, C., & Dunn, C. (1985). *Normal and disordered phonology in children*. Baltimore: University Park Press.
- Strange, W. (1992). Learning non-native phonemic contrasts: Interactions among subject, stimulus and task variables. In Y. Tohkura, E. Vatikiotis-Bateson, & Y. Sagisaka, (Eds.), *Speech Perception, Production and Linguistic Structure*. (pp. 197- 219). Tokyo: Ohmsa.
- Tees, R.C., & Werker, J.F. (1984). Perceptual flexibility: Maintenance of the recovery of the ability to discriminate non-native speech sounds. *Canadian Journal of Psychology*, 38b, 579-590.
- Templin, M. (1957). *Certain language skills in children: their development and interrelationships*. Minneapolis: University of Minnesota.
- Trubetzkoy, N. (1939). *Grundzüge der Phonologie*; translated by C. Baltaxe. 1969. *Principles of Phonology*. Berkeley: University of California Press.
- Vago, R.M. (1976). More evidence for the feature [grave]. *Linguistic Inquiry*, 7, 716-674.
- Vago, R.M. (1989). Empty consonants in the moraic phonology of Hungarian. *Acta Linguistica Hungarica: An International Journal of Linguistics*, 39 (1-4), 293-316.
- Vance, T. (1987). *An introduction to Japanese phonology*. Albany, NY: State University of New York Press.

- van der Hulst, H. (1989). Atoms of segmental structure: Components, gestures, and dependency. *Phonology*, 6 (2), 253-284.
- Vennemann, T., & Ladefoged, P. (1973). Phonetic features and phonological features. *Lingua* 32, 61-74.
- Vihman, M., Ferguson, C. & Elbert, M. (1986). Phonological development from babbling to speech: Common tendencies and individual differences. *Applied Psycholinguistics*, 7, 3-40.
- Weinreich, U. (1953). *Languages in contact: Findings and problems*. New York : Linguistic Circle of New York.
- Weiss, L. (1970). *Auditory discrimination and pronunciation of French vowel phonemes*. Unpublished doctoral dissertation, Stanford, CA. Dissertation Abstracts International, 31, 4051A.
- Werker, J., & LaLonde, C. (1988). Cross-language speech perception: Initial capabilities and developmental change. *Developmental Psychology*, 24, 672-683.
- Werker, J., & Logan, J. (1985). Cross-language evidence for three factors in speech perception. *Perception & Psychophysics*, 37, 35-44.
- Werker, J., & Polka, L. (1993). Developmental change in speech perception: New challenges and new directions. *Journal of Phonetics*, 21, 83-101.
- Werker, J., & Tees, R. (1984). Cross-language speech perception: Evidence for perceptual reorganization during the first year of life. *Infant Behavior and Development*, 7, 1866-78.
- Wexler, K., & Manzini, R. (1987). Parameters and learnability in binding theory. In T. Roeper & E. Williams. (Eds.), *Parameter Setting*, (pp. 41-76). Dordrecht: Reidel.
- White, L. (1985). Is there a "logical problem" of second language acquisition? *TESL Canada Journal*, 2 (2), 29-41.
- White, L. (1989). *Universal grammar and second language acquisition*. Amsterdam: John Benjamins.
- White, L. (1992). Long and short verb movement in second language acquisition. *Canadian Journal of Linguistics*, 37, 273-286.

- White, L. (1996). Universal grammar and second language acquisition: Current trends and new directions. In W. Ritchie & T. Bhatia, (Eds.), *Handbook of Second Language Acquisition*. (pp. 85-120). San Diego: Academic Press, Inc.
- Whitney, W.D. (1885). *Sanskrit grammar*. Cambridge: Harvard University Press.
- Yip, M. (1988). The obligatory contour principle and phonological rules: A loss of identity. *Linguistic Inquiry*, 19 (1), 65-100.
- Yip, M. (1989). Feature geometry and co-occurrence restrictions. *Phonology*, 6 (2), 349-374.
- Yip, M. (1991). Coronals, consonant clusters and the coda condition. In C. Paradis & J.F. Prunet, (Eds.), *The Special Status of Coronals: Internal and External Evidence, II: Phonetics and Phonology*. San Diego: Academic Press.