

INFORMATION TO USERS

This manuscript has been reproduced from the microfilm master. UMI films the text directly from the original or copy submitted. Thus, some thesis and dissertation copies are in typewriter face, while others may be from any type of computer printer.

The quality of this reproduction is dependent upon the quality of the copy submitted. Broken or indistinct print, colored or poor quality illustrations and photographs, print bleedthrough, substandard margins, and improper alignment can adversely affect reproduction.

In the unlikely event that the author did not send UMI a complete manuscript and there are missing pages, these will be noted. Also, if unauthorized copyright material had to be removed, a note will indicate the deletion.

Oversize materials (e.g., maps, drawings, charts) are reproduced by sectioning the original, beginning at the upper left-hand corner and continuing from left to right in equal sections with small overlaps.

Photographs included in the original manuscript have been reproduced xerographically in this copy. Higher quality 6" x 9" black and white photographic prints are available for any photographs or illustrations appearing in this copy for an additional charge. Contact UMI directly to order.

Bell & Howell Information and Learning
300 North Zeeb Road, Ann Arbor, MI 48106-1346 USA
800-521-0600

UMI[®]

NOTE TO USERS

Page(s) not included in the original manuscript are unavailable from the author or university. The manuscript was microfilmed as received.

5

This is reproduction is the best copy available

UMI

CODA CONSTRAINTS

-OPTIMIZING REPRESENTATIONS

Takako Kawasaki
Department of Linguistics
McGill University, Montreal

August 1998

A THESIS SUBMITTED TO
THE FACULTY OF GRADUATE STUDIES AND RESEARCH
IN PARTIAL FULFILLMENT OF
DOCTOR OF PHILOSOPHY

© Takako Kawasaki
August 1998



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.

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 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.

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.

0-612-54507-5

ACKNOWLEDGEMENTS

I am very glad to have an opportunity to express my thanks to my great supervisors, Heather Goad and Glyne Piggott who have spent endless hours reading my drafts, discussing problems and providing me with suggestions. Without their kindness and patience, I would not have been able to come this far. Discussions with them always helped to bring me down to earth. If I were given a next life, I would like nothing more than to be their student again (if they do not mind). I also owe much to Lydia White who gave me an opportunity to work with her. The experience in her research project is one of the invaluable experiences in my graduate career. I also would like to thank Jonathan Bobaljik and Mark Baker, for sharing their time and thoughts, as well as Lisa Travis who helped me to adapt to McGill and monitored my dissertation progress.

My appreciation also goes to the members of the McGill Dept. of Linguistics: Cindy Brown, Isabel Fonte, Joyce Garavito, Martine Kessler, Mika Kizu, John Matthews, Tomoyo Oda, Joe Pater, Ileana Paul, Phil Prevost, Miwako Uesaka, Sharon Rose, Roumyana Slavakova, O. T. Stewart, Mami Takahashi and the two secretaries, Lise and Linda. I especially thank Hidekazu Tanaka and Masanori Nakamura who have been through hard stages with me, listened to my complaints, and most of all, kept me sane from my

dissertation phobia. I wish to express my additional thanks to Madelyne Goldberg, James Lapointe and Nathaniel Stone for their editorial assistance.

I would like to extend my thanks to my friends who have supported me; Izumi Fujita, Hiroko Kawasaki, Lili and Bob Kennington, Kikuko Morinaka, Ayako Nasu, Masaru Okubo, Teiichi Ota, Nathaniel Stone and Osamu Yamamoto.

I also would like to thank the many linguists who kindly provided me with comments and directed me to sources; Jorge Guitart, Mark Hale, James Harris, John Jensen, Keren Rice, and Moira Yip. I also wish to thank Prof. Minoru Tanaka—who encouraged me to go to grad school in North America. If I had not met him, I would not have been introduced to the field of linguistics. Charles Jones—who let me know what a hard-worker I could be. Since they were few and far between, compliments from him are still carved in my mind. I cannot thank Steven Weinberger enough. He introduced me to the field of Phonology and let me feel that Linguistics is a fascinating field. Even when I could not feel confident in myself, I always have trusted what he had to say about me.

Finally, I wish to express my first appreciation to my family, Yumiko and Tomio Kawasaki, and of course, to the late Masako Kawasaki who had wished for me to become a (medical) doctor. This dissertation is dedicated to her.

TABLE OF CONTENTS

1 . INTRODUCTION	1
2 . THEORETICAL ASSUMPTIONS	9
2.1 PHONOLOGICAL FEATURES	9
2.1.1 FEATURE GEOMETRY	10
2.1.1.1 SV-Hypothesis	14
2.1.1.2 Laryngeal Node	21
2.1.2 SONORITY NODES	23
2.1.3 VOICING IN SONORANTS	24
2.2 LICENSING AND SYLLABLE STRUCTURE	30
2.2.1 SYLLABLE STRUCTURE	30
2.2.2 CODA LICENSING	31
2.2.3 GOVERNMENT PHONOLOGY	32
2.3 OPTIMALITY THEORY	34

2.3.1	CONSTRAINT INTERACTION	34
2.3.2	OPTIMALITY THEORY AND FEATURE GEOMETRY	39
2.4	THE ROLE OF [VOICE] IN THE SONORITY HIERARCHY	41
2.5	SUMMARY	48
3	. VOICING AND CODA CONSTRAINTS	54
3.1	CODA CONSTRAINTS	55
3.1.1	NOCODA: LARYNGEAL	55
3.1.1.1	Voicing Neutralization	65
3.1.2	NOCODA: PLACE	69
3.2	CODA SONORITY REQUIREMENT	71
3.2.1	CODA—SONORANT ONLY	72
3.2.2	CODA—NASAL ONLY	76
3.2.3	CODA—[ʔ]	79
3.3	LARYNGEAL ASSIMILATION	83
3.3.1	FINAL EXCEPTIONALITY	87
3.4	SUMMARY	94
4	. VOICING ASSIMILATION AND SONORITY	100
4.1	REGRESSIVE VOICING ASSIMILATION	101
4.2	POST-SONORANT VOICING	106
4.2.1	VOICING PARADOX	106
4.2.2	LICENSING CANCELLATION	107
4.2.3	LICENSING CANCELLATION AND PARASITIC LICENSING	111

4.2.4	ASYMMETRY IN POST-SONORANT VOICING	114
4.2.4.1	NoLink	114
4.2.4.2	*NC̥	118
4.2.4.3	*Complex	120
4.2.5	PRENASALIZED SEGMENTS	124
4.3	POST-SONORANT VOICING AND PRE-SONORANT VOICING	128
4.3.1	FEATURE LICENSING CONDITION	129
4.3.2	PRE-SONORANT VOICING	132
4.3.2.1	Sonority Constraint	133
4.3.2.2	Typology	138
4.3.2.3	Problem of Counting	139
4.4	SUMMARY	140
5	. YAMATO-JAPANESE	144
5.1	VOICING PARADOX IN YAMATO-JAPANESE	144
5.1.1	RICE(1993)	147
5.1.2	ITÔ, MESTER & PADGETT (1995)	149
5.2	CODA CONSTRAINTS IN JAPANESE	153
5.3	VOICING AND SONORANTIZATION	156
5.3.1	CONSTRAINT-RANKING AND CODA VOICING (YAMATO-JAPANESE)	156
5.3.2	POST-NASAL VOICING	159
5.3.3	VOICED OBSTRUENTS IN CODA	162
5.3.4	OUTPUT NEUTRALIZATION	164

5.4	[VOICE] SPECIFICATION FOR SONORANTS—AGAINST LEXICON OPTIMIZATION	166
5.5	CROSSLINGUISTIC CONSEQUENCES OF NoCODA: LARYNGEAL	170
5.6	SUMMARY	172
<u>REFERENCES</u>		<u>180</u>

I . INTRODUCTION

Languages differ in their sound patterns, but these differences are, to a large extent, systematic. One goal of Universal Grammar (Chomsky 1957, 1965) is to account for the systematic patterns which are attested across languages. Toward this end, Universal Grammar is considered to contain a set of phonological primitives such as features, and some restrictions on their combination. However, in rule-based phonology, it is assumed that rules are part of the grammar of an individual language. By their very nature, rules describe *operations*. As such, they are not well-suited to express restrictions on the ways in which segments may combine when no overt operation is involved. To account for such restrictions, Chomsky & Halle (*Sound Pattern of English* (SPE): 1968) supplemented rules with Morpheme Structure Constraints (MSCs) which define the possible morpheme shapes that a particular language allows (see also Halle 1959). Thus, in SPE, both MSCs and rules played a role in accounting for the phonological patterns observed in languages.

This dual system has many problems, one of which is the introduction of redundancy or duplication into the grammar (cf. Postal 1968). The phonological shapes which rules create and those which MSCs enforce often overlap. Secondly, as noted by Stanley (1967), there exist phonological patterns (restrictions on syllable structure for example) which cannot be expressed straightforwardly by rules and, hence, require constraints alone. Thirdly, Kisseberth (1970) points out that several rules often conspire to

satisfy a single constraint, and this functional relatedness cannot always be expressed through rule formalism. All of these problems foreshadowed the move toward a theory where rules have a minimal role to play. This trend began with the development of nonlinear phonology in the mid-1970's through the 1980's, where the move toward highly articulated representations helped to constrain the operation of rules. As representations became more elaborate, the role of the rule component was lessened in favor of constraints on representations.

Recently, many theories have placed more prominence on the role of constraints than did traditional rule-based approaches. These include Government Phonology (Kaye, Lowenstamm & Vergnaud 1985, 1990), the Theory of Constraints and Repair Strategies (Paradis 1988a, b), Declarative Phonology (Scobbie 1991, 1992), Harmonic Phonology (Goldsmith 1993), and Optimality Theory (McCarthy & Prince 1993, Prince & Smolensky 1993).¹ Perhaps the most significant advantage of approaches which shift the focus toward constraints is that they attempt to express the functional relatedness of phonological phenomena which could not be straightforwardly expressed in terms of rules. For example, vowel insertion and syllable-final consonant deletion are both strategies which languages use to avoid coda consonants. However, when they are described by two separate rules, their functional similarity cannot be detected from the rules themselves. In an approach like Optimality Theory, both phenomena are explained as consequences of a constraint which militates against the presence of coda consonants. The fact that some languages satisfy this constraint through epenthesis and others through deletion is secondary.

Among constraint-based approaches, Optimality Theory has received the most attention in the recent phonological literature. Optimality Theory differs from standard rule-based approaches as well as from other constraint-based approaches in combining two premises. One, it abandons rules and derivation altogether; two, all constraints are

considered to be universal and violable. For any given language, the set of constraints are ranked in a strict dominance hierarchy which links input (underlying representation) and output (surface representation).

As the move toward highly-articulated representations in the 1980's was accompanied by a shift toward an emphasis on constraining possible operations, we might have expected that a theory without rules like Optimality Theory would place an especially high emphasis on representations. However, this is not the case. On the contrary, in much of the literature on Optimality Theory, representational restrictions are subsumed under constraints (see Cole & Kisseberth (1994) and Pulleyblank (1994) for example). If the role of representation is subsumed under constraints, constraints must explicitly refer to the constituency and dependency relations that subsegmental structure captures. In addition, since constraints are in principle violable, the result is that *any* representational restriction can be violated. I argue that this move opens up too many possibilities. The main contribution of this thesis is to demonstrate that the combination of highly articulated representations with Optimality Theory's view of constraint violability captures several phonological phenomena in a sufficiently restrictive manner.

I will bring out the importance of subsegmental structure (feature geometry) in the theory of constraint interaction by investigating restrictions on coda position. I will focus in particular on the various restrictions which languages place on laryngeal and sonorant features in coda. I will argue that sonority and laryngeal restrictions are both due to a single constraint, one which bans a Laryngeal node in coda. I will propose further that the Spontaneous Voice (SV)² and Laryngeal nodes define sonorancy and obstruency respectively, and that, together, they make up the class of "Sonority" nodes. In addition, I will argue that the feature [voice] is dependent on both of these nodes.

The importance of this organization of features becomes particularly apparent in Chapter 5 when an account is provided for the Yamato-Japanese facts shown in (1) below. First, coda nasals and coda voiced obstruents spread voicing to the following obstruent, (1a-c). Thus, despite the widely held assumption that [voice] for sonorants is redundant and hence unspecified, nasals, together with voiced obstruents, appear to ‘spread’ voicing to the following obstruent. Second, the language also has coda sonorantization, in which coda obstruents become sonorants in coda; see (1b-d).

- (1)
- | | | | | | |
|----|---------|---|-----------|---|------------|
| a. | yom | + | te | → | yonde |
| | ‘read’ | | gerundive | | ‘reading’ |
| b. | yob | + | te | → | yonde |
| | ‘call’ | | gerundive | | ‘calling’ |
| c. | kag | + | te | → | kayde |
| | ‘smell’ | | gerundive | | ‘smelling’ |
| d. | kak | + | te | → | kayte |
| | ‘write’ | | gerundive | | ‘writing’ |

In accounting for the Japanese facts above, the structure of sonority nodes and the dominance relationship that holds between these nodes and the feature [voice] become crucial. With the particular geometry proposed, I will argue that voicing assimilation triggered by sonorants as in (1a) is a case where the [voice] feature of the SV node is parasitically licensed by the following Laryngeal node. Further, I will argue that, although they appear to be unrelated, coda sonorantization and voicing assimilation are formally processes of the same type—both are the result of a constraint which bans Laryngeal in coda.

Related to the voicing assimilation in (1), I will demonstrate that voicing assimilation is in general restricted by the prosodic relations that hold across adjacent positions. In Chapter 4, I will discuss two types of voicing assimilation that are triggered by sonorants.

all coda segments in (1) become sonorants in the output, only underlyingly voiced obstruents and underlying nasals trigger voicing assimilation; underlyingly voiceless obstruents do not trigger this process. Comparing underlying voiced obstruents with voiceless obstruents ((1b, c) and (1d), respectively), we can conclude that only segments which have [voice] in the input trigger voicing. If [voice] is not specified for nasals (see (1a)), we cannot account for why nasals trigger voicing assimilation together with voiced obstruents. To account for this fact, I will argue that [voice] must be specified for sonorants in the input (see Section 5.4).

The thesis is organized as follows. In Chapter Two, I will summarize my theoretical assumptions regarding subsegmental structure, as well as the basic premises of Optimality Theory. Chapter Three provides an overview of coda constraints, focusing mainly on constraints which affect the Laryngeal node. In this chapter, I will show how various ways of satisfying a coda constraint on Laryngeal yields languages where coda obstruents become sonorants, languages where coda laryngeals are neutralized, and languages where codas assimilate the Laryngeal specification of the following onset. Chapter Four deals with voicing assimilation, with the main focus being on assimilation triggered by sonorants. It addresses two problems concerning this process: the redundancy of the feature [voice] for sonorants, and the directionality of assimilation. In accounting for why progressive and regressive voicing assimilations pattern differently, I will argue that headedness relations across syllables are required. In particular, as mentioned earlier, I will demonstrate how the notion of ‘head’ in Government Phonology accounts for certain directional asymmetries. In Chapter Five, I will investigate phonological restrictions in Yamato-Japanese which are correlated with coda constraints. I will show how the proposals put forward in earlier chapters can capture both voicing assimilation and sonorantization in Yamato-Japanese. Subsequently, I will show that the proposed feature geometry, together with constraints on

coda, predicts the existence of five different language types, all of which are attested: 1) languages where coda obstruents become sonorants (Hausa), 2) languages where laryngeal features are neutralized in coda (German, Maidu), 3) languages where coda obstruents become glottal stop (Kiowa), 4) languages where coda voiced obstruents assimilate in laryngeal specification to the following onset (Ancient Greek), and 5) languages where coda obstruents are unaffected (English).

—NOTES—

¹ This approach also holds true of works within the principles and parameters framework of Chomsky (1981), e.g. Piggott & Singh (1985), Itô (1986), Singh (1987), Archangeli & Pulleyblank (1994).

² Motivation for an SV node is provided in Avery & Rice (1989, 1991) and Piggott (1992). See further Chapter 2.

2 . THEORETICAL ASSUMPTIONS

In this chapter, I will introduce the ideas on which my analyses will be based. First, I will review some approaches to the organization of phonological features. In Section 2.1, I will discuss the feature geometry that I assume, focusing on sonority features and laryngeal features. Secondly, in Section 2.3, I will summarize the basic premises of Optimality Theory, and explain the consequences of combining this theory with featural organization. Lastly, I will show how the Sonority scale can be encoded in terms of the proposed structure.

2.1 Phonological Features

It is widely accepted that segments are not phonological atoms but consist of smaller units, *features*. The issue of whether features are unary or binary has received much attention within feature theory. From Jakobson (1941) through to the early 1980s, it had generally been assumed that all features are binary (+ or -). However, another view has emerged since the development of feature geometry. In early models, non-terminal features (organizing nodes) were assumed to be monovalent and terminal features bivalent. For example, in

Sagey's (1986) model, non-terminal place features are correlated with the articulator used to produce a sound. They are thus, by definition, monovalent. Since then, there has been a move toward the view that all features and nodes are monovalent (e.g., Rice & Avery 1991). This view has been motivated not only by considerations of parsimony, but also on empirical grounds. In phonological processes, there is a significant asymmetry observed between the two values of most features. Most phonological generalizations make reference to either the positive or negative value of a feature while few refer to both values. Some phonologists have defended monovalency for individual features or groups of features (Laryngeal features (Itô & Mester 1986; Lombardi 1991), [nasal] (Piggott 1992), Height features (Goad 1993)); others have proposed that all features have this property (Schane, 1984a, b; Anderson & Ewen 1987; van der Hulst 1989; Rice & Avery 1991).¹ I follow the latter position and assume that all features and nodes in the geometry are privative.

The study of phonological features and how they are organized has advanced significantly over the last decade. There are two main tenets: 1) feature arrangement and 2) underspecification. In this section, I will discuss each of these in turn.

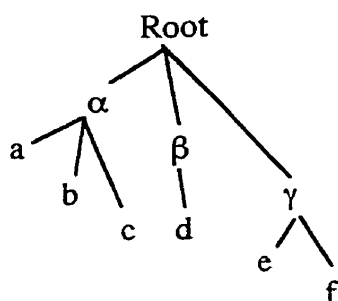
2.1.1 Feature Geometry

Many works have been devoted to demonstrating that features are hierarchically organized, developing a model called *Feature Geometry* (e.g., Mascaró 1983; Mohanan 1983; Clements 1985; Mascaró 1986; Sagey 1986). Even at an earlier period when it was believed that segments consist of bundles of features (*feature matrices*) with no internal structure, it had been observed that certain features consistently behave as a group in assimilation rules.

For example, many phonologists working within the SPE framework (Chomsky & Halle 1968) made use of abbreviations such as [α Place] in describing the cases in which all place features pattern together in phonological processes.

Feature Geometry expresses such feature grouping by means of structural constituency—functionally-related features are grouped under a single organizing node.

(1) Feature Geometry Model



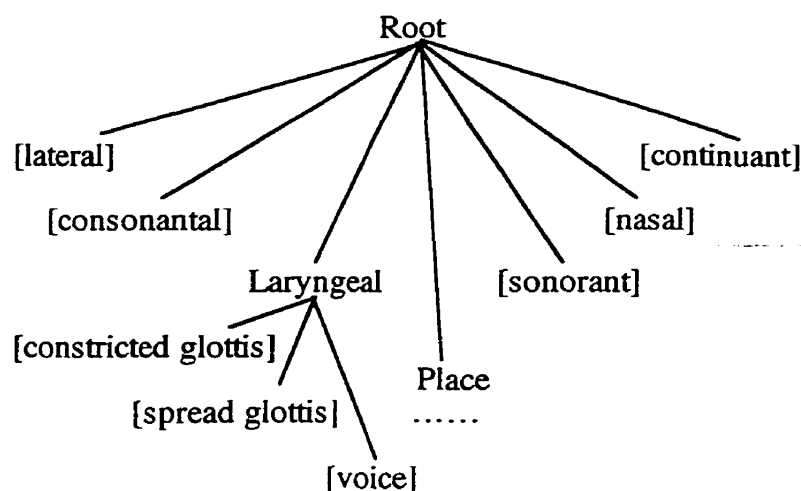
In (1), *a-f* are terminal features which are dominated by the organizing nodes α , β and γ . These nodes are in turn dominated by a higher node, Root, which organizes all nodes and features in the configuration. In this model, the features *a*, *b* and *c* are expected to pattern as a unit; that is, if node α spreads or deletes, then all features dominated by this node will be affected.

A basic premise of Feature Geometry is that each feature and its organizing node(s) are in a dominance relation as (1) illustrates. Presence of either feature *e* or feature *f* entails the presence of its mother node, γ . That is, *e* and *f* must link to Root through γ .

Since the first models of feature geometry were proposed, many modifications have been made. As the main focus of this thesis is the organization of sonority and voicing

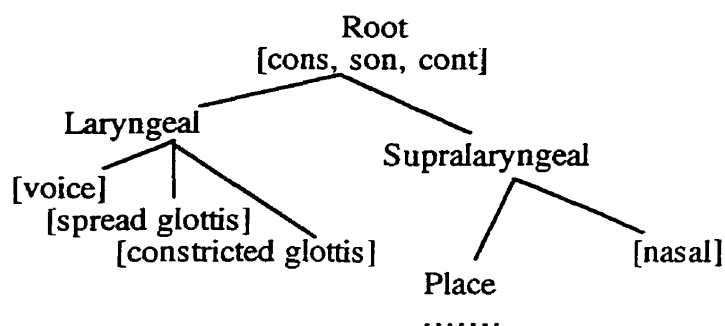
features, I will discuss the modifications proposed for these features. I will start with the geometry put forth by Archangeli & Pulleyblank (1989, 1994) shown in (2) and will discuss two issues: 1) the location of [sonorant], 2) the Laryngeal node and its dependents.

(2) Archangeli & Pulleyblank's (1994) Feature Geometry²



In the geometry in (2), the feature [sonorant] is a direct dependent of the Root node. The location of this feature has been subject to modification, as have other parts of the geometry. Consider Schein and Steriade's geometry in (3).

(3) Schein and Steriade's (1986) Feature Geometry

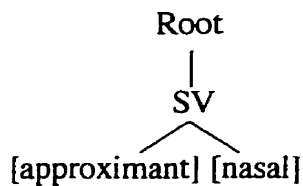


Notice that in this geometry, sonorant-obstruent specification is part of the Root node, a hypothesis which is also adopted by McCarthy (1988) and others. This proposal is justified by the claim that the 'major class features [sonorant] and [consonantal] differ from all other features in one important respect: they arguably never spread, delink, or exhibit OCP effects ...' (McCarthy 1988:97). If a feature defines the Root node, operations cannot act on it independently.

However, the claim that major class features define the Root node has more recently been challenged by Rice & Avery (1991). Rice & Avery recognize that there are phenomena which indicate that sonorancy changes without affecting place of articulation. To adequately describe such phenomena, they have proposed the SV-Hypothesis (Rice & Avery 1989; Piggott 1992) which holds that sonorants bear an SV node (Sonorant Voice for Rice & Avery and Spontaneous Voice for Piggott) which, itself, dominates other features.

Rice & Avery (1989) and Piggott (1992) propose that voicing for sonorants is encoded by SV, while voiced obstruents bear the feature [voice]. The SV node organizes the features [nasal] and [approximant] (in Rice & Avery's version, [lateral] is a dependent of SV instead of [approximant]). The structure of the SV node that I adopt is given in (4) below. I assume that this node is present in all sonorants, both consonants and vowels.

(4) Structure of SV



The SV-Hypothesis is overwhelmingly supported by cross-linguistic evidence from nasal harmony (Piggott 1992a, b), nasal-liquid alternations (Rice & Avery 1991; Rice 1993) and desonorantization (Rice & Avery 1991). As I will adopt the SV-Hypothesis, I summarize some of this evidence below.

Piggott (1992a, b) observes that in languages such as Southern Barasano and Guarani (which he calls Type B languages), the feature [nasal] spreads only to sonorants (vowels, glides and liquids). Obstruents are skipped and not nasalized by this process. Some illustrative data are shown in (5) and (6) below.

(5) Southern Barasano
(Smith & Smith 1971)

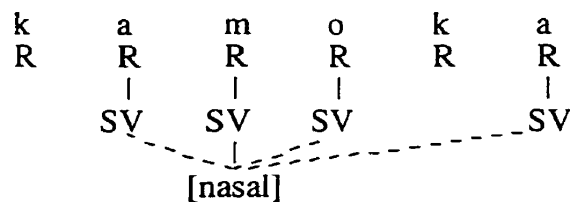
- | | |
|-------------|----------|
| a. māhāŋĩ | ‘comer’ |
| b. ŋāmōnōnĩ | ‘year’ |
| c. māsā | ‘people’ |
| d. ŋūkā | ‘drink’ |
| e. wāfĩ | ‘demon’ |

(6) Guaraní
(Piggott (1992),
originally from Rivas (1974))

- | | |
|---------|-------------|
| a. tūpā | ‘god’ |
| b. pĩrĩ | ‘to shiver’ |
| c. mēnā | ‘husband’ |
| d. nūpā | ‘to beat’ |
| e. māʔē | ‘to see’ |

Under the geometry where [sonorant] is a Root feature, (see (3) for example), there cannot be a structural explanation for why nasal harmony only targets sonorants. Instead, in order to account for such cases, a feature co-occurrence constraint, *[nasal, -sonorant] (Archangeli & Pulleyblank 1989), is needed which prohibits a segment from bearing both [nasal] and [—sonorant]. However, the co-occurrence constraint does not provide a non-arbitrary solution for why obstruents can be skipped by nasal harmony. With the geometry in (4), on the other hand, nasal harmony in (5) and (6) can be expressed as [nasal] spreading to SV nodes as (7) illustrates.

(7) Nasal Harmony in SV Geometry



Thus, the SV-Hypothesis allows us to structurally capture the group of targeted segments, namely SV-bearing segments. Since obstruents do not bear a SV node, they can be transparent to this type of nasal harmony.³

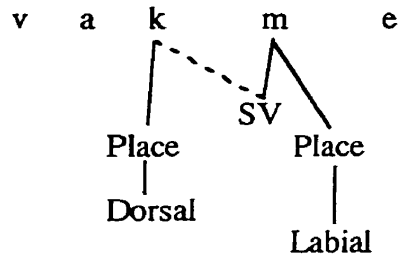
Rice and Avery (1991) provide evidence of a different nature for the SV-Hypothesis. They focus on liquid-nasal alternations. In Sanskrit, for example, stops assimilate to following nasals and liquids. As the examples in (8) show, changes in [sonorant], [nasal] and [approximant] do not affect place specification (see (8b) and (8c)).

(8) Sanskrit (Rice & Avery 1991:113, originally from Whitney 1889)

- | | | |
|------------------|---|---------------|
| a. tat namas | → | tan namas |
| b. vak me | → | vaṇ me |
| c. triṣṭup nunam | → | triṣṭum nunam |
| d. tat labhate | → | tal labhate |
| e. út luptam | → | úl luptam |

Again, this set of facts cannot be straightforwardly accounted for by the geometry in (3) where [sonorant], as part of the Root node, dominates Place. In the geometry in (3), the transfer of [sonorant] should be accompanied by the transfer of all dependent nodes and features including Place. On the other hand, under the SV hypothesis in (4), the Sanskrit facts can be expressed as spreading of the SV node as shown in (9).

(9) Sanskrit SV Assimilation⁴



Since the SV node is independent of Place, the place specification of the target segment is not affected by this process.

With respect to SV structure, I will assume that nasal is the default interpretation of a bare SV consonant, following Rice & Avery (1991) and Rice (1993) (see Chapters 4 and 5). They propose that in the unmarked case, the feature [nasal] is not specified. Their proposal of a bare SV as nasal is consistent with Kean's (1975) observation that nasals are the least marked sonorant consonants. They also support this assumption through assimilation patterns. In assimilations between liquids and nasals, nasals are usually targets. See the examples from Klamath and Ponapean below.

(10) Klamath (Barker 1964)

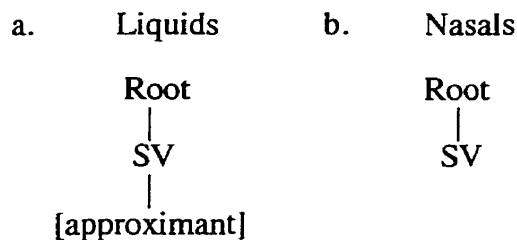
- | | | | | |
|----|-----------|---|-----------|----------------------------|
| a. | honlina | → | hollina | 'flies along the bank' |
| b. | w'inl'ga | → | w'illga | 'lies down on the stomach' |
| c. | pečqnl'ga | → | pečqallga | 'puts a foot down through' |

(11) Ponapean (Rehg & Sohl 1981)

- | | | | | |
|----|-----------|---|-----------|---------------------|
| a. | nan-leg | → | nalleng | 'heaven' |
| b. | Kepinle | → | kepille | 'a place name' |
| c. | pan ligan | → | pallingan | 'will be beautiful' |

In addition to the languages above, Toba Batak has similar assimilations. According to Hayes (1986), a coronal nasal assimilates to the following liquids while liquids do not assimilate when followed by nasals (see Hayes 1986:479). This asymmetry in assimilation can be interpreted as an indication that nasals do not have any dependent under the SV node while liquids do. Consistent with this, the SV structures for liquids and nasals that I adopt are in (12).

(12) Representations for Sonorant Consonants⁵



If the feature [nasal] is not present for nasals in the unmarked case, when does it play a role in languages? As there are clearly languages where [nasal] must be present under SV, the theory must provide two options: 1) [nasal] is not present—unmarked, 2) [nasal] is present—marked. I suggest that presence of this feature is based on positive evidence which can come in two forms: either from the presence of particular contrasts, or from phonological processes. When a language has a contrast between nasal and oral vowels,

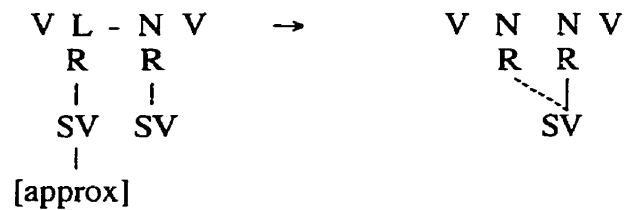
then the language requires the presence of the feature [nasal]. If a language does not have such a contrast, but it has a process such as nasal harmony which reveals the presence of the feature, [nasal] must also be projected. If there are no segments which contrast only for the feature [nasal] nor processes which reveal the presence of [nasal], I propose that this feature is absent from the language. Importantly, the presence of nasal segments in a language is not sufficient to trigger the presence of [nasal]; nasals will in the unmarked case be represented by a bare SV node. To illustrate, let us compare two hypothetical patterns in (13) below where V stands for vowels, L for liquids, and N for nasals. Underlined segments reflect the output of assimilation.

(13)

- a. V L - N V → V n n V
b. V L - N V → V n n V

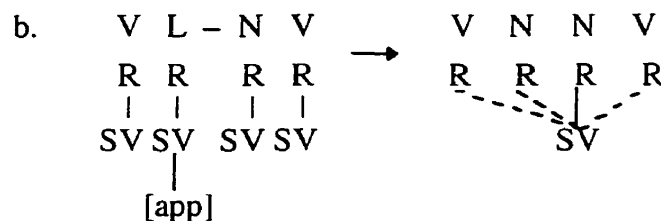
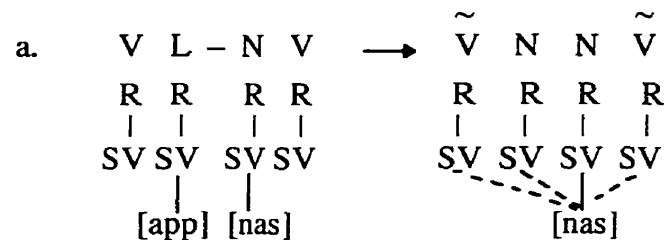
If a child encounters data of the type in (13a), s/he will not have to adopt the feature [nasal]. The liquid-nasal alternation can be captured through SV spreading as in (14).

(14)



Since the feature [nasal] is not necessary to express the type of process, data of the type in (13a) will not lead the child to project [nasal]. On the other hand, the projection of [nasal] will be triggered by data like (13b). To account for the nasal harmony in (13b), [nasal] is a necessary feature of the underlying nasal consonant. See (15).

(15)



As (15a) shows, a pattern like (13b) can best be explained as spreading of the feature [nasal]. If this process were instead treated as SV spreading, the vowel would not become nasalized as illustrated in (15b). Therefore, a pattern like (13b), in particular, the difference

between V and ∇ will lead the child to adopt [nasal]. In summary, the child projects the feature [nasal] when s/he detects minimal contrasts between nasal and oral segments or when s/he encounters phonological phenomena which cannot be expressed without the feature nasal like in (15a).

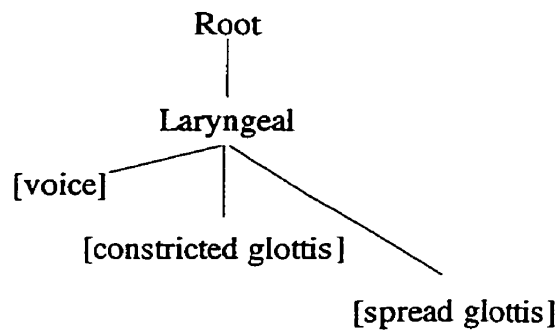
As mentioned earlier, I adopt the view that SV expresses sonorancy; that is, it replaces the feature [sonorant]. So far, we have seen motivation for positing the SV node and we have briefly discussed its dependents, [approximant] and [nasal]. Now I will turn to the organization of the Laryngeal node which defines obstruency. We will return to the organization of SV features shortly.

2.1.1.2

Laryngeal Node

It is widely accepted that Laryngeal is a node which organizes features that refer to states of the glottis: [voice], [constricted glottis] ([CG]) and [spread glottis] ([SG])⁶ (see (16)). [CG] is a feature which identifies glottalized segments and [SG] identifies aspirated segments (see Halle and Stevens (1971)).

(16) Structure of Laryngeal Node



Sub-laryngeal features behave as a group when contrasts are neutralized in coda, either as a result of deletion or of assimilation. For example, in Klamath (Barker 1963, 1964, Lombardi 1991) which has three-way laryngeal contrasts in obstruents (plain, aspirated, glottalized), all laryngeal contrasts are neutralized syllable-finally as (17) shows.

(17) Klamath Laryngeal Neutralization (Barker 1963)

- | | | |
|----|-------------------------------------|--------------------------|
| a. | /mp ^h et'/ | 'float' |
| | mp ^h et'i:qi | 'floats up' |
| | mp ^h etplanča | 'floats downstream' |
| b. | /p ^h eč ^h / | 'foot' |
| | p ^h eč ^h i:qi | 'puts a foot into water' |
| | p ^h ečk'wa | 'puts a foot across' |

In (17), both [SG] and [CG] are lost in coda. The loss of laryngeal features in (17) can be unified as loss of the Laryngeal node in coda.

In Ancient Greek, a coda assimilates to the Laryngeal specification of the following onset when the onset is Coronal obstruent. Examples are provided in (18).

(18) Ancient Greek (Steriade 1982:231)

a. kleb-dēn	‘stealthy’	:	e-klap-ēn	‘I was cheated’
b. pleg-dēn	‘entwined’	:	plekō	‘to plait’
c. tet ^h lip-tai	‘has been squeezed’	:	t ^h libō	‘to squeeze’
d. lek-teos	‘to be counted’	:	legō	‘to count’
e. strep-tos	‘turned’	:	strep ^h ō	‘to turn’
f. e-dok ^h -t ^h ē	‘it seemed’	:	dok-e-ō	‘to count’

As the left column of data shows, the Laryngeal specification of the coda obstruent is identical to that of the following onset. For example, /p/ in /klap/ (see (18a)) becomes [b] when followed by the voiced obstruent, /d/; /g/ in /leg/ (see (18d)) becomes voiceless when followed by the voiceless obstruent, /t/; and /k/ becomes aspirated [k^h] when the following coronal is aspirated (see (18f)). Importantly, these data show that both Laryngeal features which are contrastive in the language, [voice] and [SG], are assimilated. Therefore, spreading of the Laryngeal node unifies this process.

To summarize, both the neutralization and assimilation facts support a configuration where there is a Laryngeal node which dominates [CG], [SG] and [voice].

2.1.2 Sonority Nodes

While sonorancy is indicated by the presence of the SV node, I propose that the Laryngeal node is correlated with obstruency (Kawasaki 1995, 1996). These two nodes, Laryngeal

and SV, are thus functionally similar in that they characterize the contrast in sonority. I will refer to these nodes, which define the sonority of segments, as Sonority Nodes; see (19).

(19) Sonority Nodes

SV and Laryngeal are Sonority Nodes which define the sonority of a segment (i.e. sonorant vs. obstruent respectively).

If Laryngeal specifies obstruency and the Laryngeal node dominates the feature [voice], the following questions arise. How is voicing for sonorants specified? Since sonorants are intrinsically voiced, can they bear Laryngeal [voice]? Both Chomsky & Halle (1968) and Ladefoged (1982) recognize that voicing for sonorants and voicing for obstruents are fundamentally different. There are also some phonological phenomena (e.g., *Rendaku* in Japanese), in which voiced obstruents do not pattern together with sonorants, regardless of the fact that both types of segments are phonetically voiced. In order to capture both the phonetic and phonological facts with a single structure, I adopt the dual-dependency of [voice] which was first proposed by Piggott (1994) and further elaborated on in Kawasaki (1995, 1996). Under this proposal, the feature [voice] is dependent on both the Laryngeal and SV nodes (see (23) below). In the following section, I will discuss the motivation for this hypothesis.

2.1.3 Voicing in Sonorants

In many languages, when a nasal-final prefix is attached to a stem, a stem-initial voiceless obstruent becomes voiced (e.g. Kpelle (Welmers 1973), Kikuyu (Armstrong 1967), Ndali

(Vail 1972), Terena (Bendor-Samuel 1960), Maukaka of Mwaaluu (Tourville 1991)).

Examples from Kpelle and Ndali are provided in (20) and (21), respectively.

(20) Kpelle (Sagey 1986, originally from Welmers 1973)

- | | | | | |
|----|----------|---|------------------------|------------|
| a. | /N-polu/ | → | [m ^h bolu] | ‘my back’ |
| b. | /N-tia/ | → | [n ^h dia] | ‘my taboo’ |
| c. | /N-kpiŋ/ | → | [m ^h ŋgbiŋ] | ‘myself’ |
| d. | /N-fela/ | → | [m ^h vela] | ‘my wages’ |
| e. | /N-sua/ | → | [n ^h jua] | ‘my nose’ |

(21) Ndali (Vail 1972)

- | | | | | |
|----|---------------|---|-------------------------------------|--------------------|
| a. | /iN + puno/ | → | [i ^m buno] | ‘nose’ |
| b. | /iN + puunda/ | → | [i ^m buunda] | ‘horse’ |
| c. | /iN + tongi/ | → | [i ⁿ do ⁿ gi] | ‘lump in porridge’ |
| d. | /iN + tunye/ | → | [i ⁿ dunye] | ‘banana’ |
| e. | /iN + kunda/ | → | [i ⁿ gunda] | ‘dove’ |
| f. | /iN + kwero/ | → | [i ⁿ gwero] | ‘spear’ |

In (20) and (21), [voice] appears to be spreading from the nasal to the stem-initial obstruents. In contrast to these languages, the phonological inertness of [voice] for sonorants—including nasals—is well-documented in the literature (see Kiparsky 1982, 1985, Itô & Mester 1986). Let us review in this regard the well-known Yamato-Japanese *Rendaku* examples in (22).

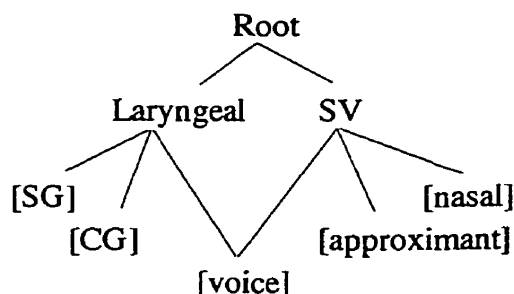
(22) Rendaku

a. ude + tokei 'wrist' 'watch'	→ udedokei 'wrist watch'
b. te + kufi 'hand' 'comb'	→ tegufi 'hand comb'
c. te + saguri 'hand' 'search'	→ tesaguri *tezaguri 'grobe'
d. ai + kagi 'match' 'key'	→ aikagi *aigagi 'spare key'
e. mizu + kame 'water' 'jar'	→ mizugame 'water jar'
f. it̥ʃigo + kari 'strawberry' 'hunting'	→ it̥ʃigogari 'strawberry picking'
g. te + sawari 'hand' 'touch'	→ tezawari 'touch'

Rendaku is a process which voices the initial segment of the second member of a Yamato-Japanese compound (22a,b). However, when the second member of a compound already contains a voiced obstruent, *Rendaku* is blocked (22c, d). The presence of sonorant consonants does not block the application of *Rendaku* (22e-g). Thus, although sonorants are phonetically voiced, they do not behave like [voice]-bearing segments with respect to *Rendaku*.

Because of facts like *Rendaku*, it has been argued that sonorants are unspecified for [voice]. However, as we saw in (20) and (21), there exist many languages where [voice] for sonorants appears to play a role in the phonology.⁷ To resolve this paradox, Piggott (1994) proposes a modification to the feature geometry which allows [voice] the option sketched in (23).⁸

(23) Dual-Dependency of [voice]



As (23) shows, [voice] is dominated by both Laryngeal and SV. The implication is that sonorants have [voice] under the SV node and obstruents have [voice] under the Laryngeal node. Piggott (1994) suggests that for the post-nasal voicing cases in (20) and (21), it is [voice] under the SV node in nasals which spreads to the Laryngeal node of the stem-initial obstruents.

The dual-dependency of [voice] not only accounts for post-nasal voicing, which will be discussed in detail in Chapter 4, but it can be extended to other facts as well. Consider Havana Spanish (espontáneo) in (24) where liquids lose sonorancy in coda. In Havana Spanish, coda liquids become voiced obstruents when followed by stops (Harris 1986). The following examples illustrate this.

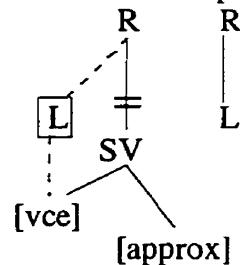
(24) Havana Spanish (Harris 1986)

- | | |
|----------------------|---------------|
| a. g <u>o</u> rdo | → go[dd]o |
| b. pu <u>r</u> ga | → pu[gg]a |
| c. e <u>l</u> te | → e[d t]e |
| d. se <u>r</u> pobre | → se[b p]obre |
| e. ta <u>l</u> nata | → ta[d n]ata |

When we consider the data in (24c-d), it becomes clear that this is not a gemination process since the voicing specifications are not the same between the two segments in the output. Instead, these forms manifest the relationship that holds between sonorants and voiced obstruents. With the structure in (23), this relationship can be straightforwardly expressed.⁹ In (25), coda liquids lose their SV specifications. While one SV dependent, [approximant], must be lost through delinking of the SV node, the other SV dependent, [voice], can be saved through attachment to another mother node, Laryngeal. Thus, according to the adopted geometry, the phenomenon in (24) can be treated as a change in the mother node of [voice]: SV becomes Laryngeal.

(25) Havana Spanish Obstruentization¹⁰

s e r p o b r e (→seb pobre)



The proposal in (23) has several implications. First, this geometry suggests that voicing for sonorants and voicing for obstruents are the same, yet different. They are the same since voicing in both types of segments is specified by a single feature [voice]. However, the feature [voice] is dominated by different nodes—SV for sonorants and Laryngeal for obstruents. Both the similarity and the difference are phonetically motivated. Voicing for sonorants and voicing for obstruents are both produced by vibration of the vocal

cords, and this is expressed by the single feature [voice]. However, the vocal cavity configurations for sonorants and voiced obstruents are different (Chomsky & Halle 1968, Ladefoged 1982). This difference is expressed by the difference in mother nodes; the SV node designates a vocal cavity configuration which makes spontaneous voicing possible (sonorants) while the Laryngeal node defines a vocal cavity configuration where *spontaneous* voicing is impossible (obstruents). Thus, although voicing itself is specified by the same feature, the combinations of the feature and its mother nodes express the phonetic differences.

The dual-dependency of [voice] has another implication when combined with the proposal in (19) where SV and Laryngeal are defined as sonority nodes. The geometry in (23) places [voice] under both sonority nodes. This feature is therefore expected to play a role in determining the sonority value of a segment. This will be addressed in Section 2.4.

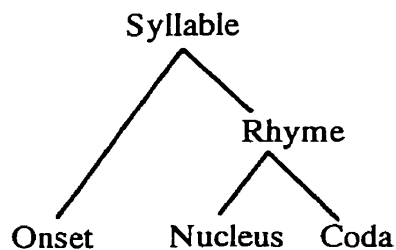
So far, I have discussed the subsegmental organization of the SV and Laryngeal nodes. By incorporating feature geometry into the theory of constraint-interaction, which will be introduced in Section 2.3, I will demonstrate that hierarchical relations still play an important role in phonology. I will argue that this is true not only at the level of the segment, but at the level of higher prosodic structure as well. In the next subsection, I will discuss the organization of prosodic categories and the relations between syllable positions that I adopt.

2.2 Licensing and Syllable Structure

2.2.1 Syllable Structure

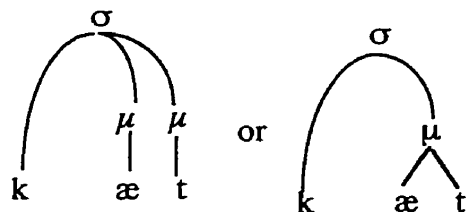
In the standard theory, segments are dominated by subsyllabic constituents onset, rhyme, nucleus and coda, which are organized as in (26).

(26) Syllable Structure in the Standard Theory



More recently, another syllable-internal element, the *mora* (μ), has been proposed, thereby challenging the traditional syllable constituents as onset, nucleus, and coda; see (27) (Hyman 1984, 1985, McCarthy & Prince 1986, Hayes 1989).¹¹ In Moraic Theory, the mora serves to indicate both weight and position.¹² The number of morae a syllable has determines its weight (i.e., whether it is heavy (bimoraic) or light (monomoraic)). Whether or not coda consonants are moraic varies across languages; the contrast is illustrated below.

(27) Syllable Structure in Moraic Theory



In the theories proposed by Hyman (1984, 1985), McCarthy & Prince (1986), and Hayes (1989), onset, nucleus and coda are not constituents. Therefore, the means to refer to coda position will be different from in Onset-Rhyme Theory. In Moraic Theory, a coda consonant can be described as a consonant dominated by a mora. Although I will refer to the post-nucleus position as a “coda”, in the analyses which will be presented, nothing rests on this label, and I do not wish to take a position on the correctness of either conception of syllable structure in this thesis. In the following section, we will discuss the issue of licensing in this coda position.

2.2.2 Coda Licensing

Syllables are dominated by higher prosodic categories, minimally the Foot and the Prosodic Word. This hierarchy plays an important role in the theory of prosodic licensing proposed by Itô (1986). Itô proposes that all phonological entities must be licensed by higher prosodic structure in order for them to be parsed (phonetically realized). When it comes to terminal syllable positions, their “licensing power” is not equal. As is widely recognized, segments which can appear in coda are more limited than those which can appear in onset. This is

because the licensing ability of codas is more limited than that of onsets. The restrictions that languages impose on codas have been expressed in terms of coda conditions (or constraints) (see Itô 1986, Lombardi 1991, among others).

The main focus of this thesis is to explore the consequences of such constraints on licensing. Languages resolve coda violations differently, either by delinking nodes which are the target of some restriction (neutralization), or by multiply linking the targeted node to the following onset (assimilation). For example, among languages where codas are not allowed to license a Laryngeal node, German chooses to delink Laryngeal in coda while Dutch chooses to have the coda share Laryngeal with the following onset. In Chapter 3, I will demonstrate how the proposed feature geometry accounts for the different types of resolutions of coda restrictions.

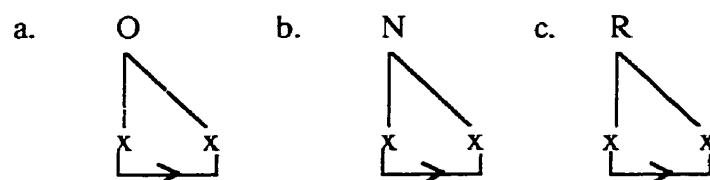
2.2.3 Government Phonology

In addition to exploring the licensing power of codas, I will also argue that the onset-coda relation plays an important role in phonological alternations (see Chapter 4). Languages display asymmetries between coda-to-onset (progressive) assimilations and onset-to-coda (regressive) assimilations. To account for the directional asymmetries observed, I adopt some of the basic concepts of Government Phonology (Kaye, Lowenstamm & Vergnaud 1985, 1990, Harris 1990, Kaye 1990, 1994).

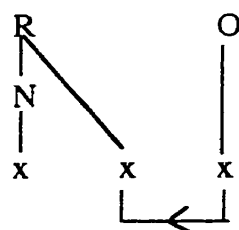
Phonological government is an asymmetric relation that holds between elements in certain prosodic configurations; it is strictly local and binary. Governing relations are of two types: constituent government applies within syllable constituents and interconstituent government applies between constituents. Government relations are unidirectional: in

constituent government, the head position is initial while in interconstituent government, the head is final. (28) and (29) illustrate constituent government and interconstituent government, where O, R, N refer to Onset, Rhyme and Nucleus, respectively.

(28) Constituent Government



(29) Interconstituent Government



Here, we are concerned with interconstituent government where the government relationship is from right-to-left; i.e., the coda (dependent) is governed by the following onset (head) and the reverse relation does not hold. I will propose that, because of this strictly unidirectional relationship, licensing is also unidirectional. This notion of head will thereby enable us to account for the asymmetrical nature of assimilation which will be discussed in Chapters 4 and 5.

In the following sections, I will provide an overview of Optimality Theory, which is the framework in which my analyses are couched. I will argue that the phenomena to be investigated are best explained within a theory of constraint interaction which assumes subsegmental structure.

2.3 Optimality Theory

As discussed in Chapter 1, in rule-based frameworks which have been adopted since Chomsky and Halle (1968), the mapping between underlying and surface representation is achieved through a series of ordered language-specific rules, the application of which is sometimes governed by well-formedness constraints. Optimality Theory, on the other hand, is a theory of constraint interaction which denies the existence of rules and derivations altogether (Prince & Smolensky 1993). The mapping of underlying representation to surface representation is regulated by a set of universal constraints which are variably ranked across languages.

2.3.1 Constraint Interaction

In Optimality Theory, a grammar consists of two components: GEN and EVAL. GEN is a function which produces a number of candidate outputs (potential surface representations) for a given input (underlying representation). Outputs generated by GEN thus include many forms which are not realized on the surface in a particular language. The various outputs are fed into EVAL which consists of a universal set of constraints, rank-ordered across languages. Since all constraints are claimed to be universal, they must also be violable. Within a given language, higher ranked constraints have absolute priority over lower ranked ones. The optimal output is the candidate which best satisfies the given constraint ranking, i.e. that which violates the fewest highly ranked constraints.

The outputs that GEN produces may differ from the input in many respects; for example, at the segmental level, any feature or node can be inserted or deleted, thereby

obscuring the identity between input and output. Constraints on input-output identity are expressed in terms of a relation called Correspondence which is defined in (30) below.

(30) Correspondence (McCarthy & Prince 1994)

Given two strings S_1 and S_2 , related to one another as reduplicant/base, output/input, etc. correspondence is a function f from any subset of elements of S_2 to S_1 . Any element α of S_1 and any element β of S_2 are correspondents of one another if α is the image of β under correspondence; that is, $\alpha = f(\beta)$.

Each candidate input-output string is assessed for its identity in terms of correspondence. The constraints which are designed to ensure input-output identity are called *Faithfulness* constraints. Two subfamilies of these constraints are described in (31).

(31) Faithfulness Constraints¹³

MAX Every element of S_1 has a correspondent in S_2 .

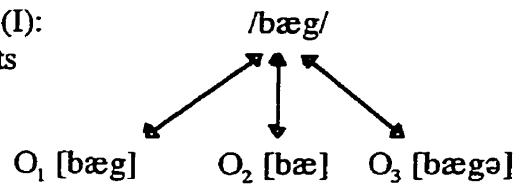
DEP Every element of S_2 has a correspondent in S_1 .

MAX is a family of constraints which ensures that every element in the input has a corresponding element in the output. In other words, MAX bans deletion.¹⁴ In contrast, DEP ensures that every element in the output has a correspondent in the input; i.e., DEP is a constraint on insertion.

To exemplify how MAX and DEP strive for faithfulness between input and output, let us consider three candidate outputs for the input /bæɡ/.

(32) Input-Output Correspondence

Input (I):
Outputs



In the illustration in (32), correspondence relations between the input and outputs O_1 , O_2 and O_3 are expressed by arrows. Each pair which is linked by an arrow is subject to evaluation. The pair I and O_1 are in perfect I-O (Input-Output) correspondence, since every element in I has a correspondent in O_1 which is its exact image. However, in O_2 , the last segment /g/ in I does not have a correspondent (deletion). In O_3 , the final vowel [ə] does not have a correspondent in I (insertion). The second and third candidates each violate one of the *Faithfulness* constraints. (33) identifies the particular Faithfulness constraint that is violated by the candidates in (32).

(33)

input: /bæɡ/

		MAX	DEP
O_1	[bæɡ]	✓	✓
O_2	[bæ]	*	✓
O_3	[bæɡə]	✓	*

Although O_1 satisfies both Faithfulness constraints, it is not necessarily selected as the optimal candidate. This is because there is a tension between input-output faithfulness and structural well-formedness. The latter prefers candidates which are structurally unmarked at the expense of violating faithfulness.¹⁵ For example, in a language which does not allow coda consonants (e.g., Senufo (Mills 1984)), O_1 cannot surface as optimal. In such a language, a markedness constraint prohibiting codas (NoCODA) dominates one of the faithfulness constraints. The definition of NoCODA is provided in (34).

(34)NoCODA

*C]_σ, i.e. codas are prohibited (Prince and Smolensky 1993)

The tableaux in (35) illustrate how different rankings between NoCODA and Faithfulness will select different outputs as optimal.

(35)input: /bæɡ/¹⁶

a.		DEP	NoCODA	MAX
	O_1 bæɡ		*	
→	O_2 bæ			*
	O_3 bæɡə	*		

b.		MAX	NoCODA	DEP
	O_1 bæɡ		*	
	O_2 bæ	*		
→	O_3 bæɡə			*

In the tableau in (35a), DEP (which bans insertion) and NoCODA (which prohibits codas) are ranked higher than MAX (which prohibits deletion). Therefore, candidate O_2 in which the coda segment is deleted is selected as optimal. In the tableau in (35b), MAX and NoCODA are ranked higher than DEP. Therefore, candidate O_3 in which a vowel is inserted to make the

last consonant an onset is selected. If both MAX and DEP outrank INOCODA, O_1 which is the perfect correspondent of the input, will be selected. As we have seen, reranking of NOCODA and the Faithfulness constraints yields three types of outputs. Thus, in Optimality Theory, cross-linguistic variation is mainly due to differences in constraint-ranking.

As briefly mentioned in Chapter 1, one important advantage of a constraint-based approach such as Optimality Theory is that it allows us to express relatedness across phonological phenomena which are attested in many languages, but that manifest themselves in slightly different ways. In a rule-based approach, such differences across languages must be accounted for by rules which are often very different from one another. Thus, the fact that several rules may conspire to satisfy the same constraint is not reflected in the formalism (Kisseberth 1970). Mohanan (1993) has raised this theoretical concern in the following way.

(36) Mohanan's Central Questions (Mohanani 1993:66)

How do we capture the universality of those patterns of distribution and alternation which

- (i) appear repeatedly in human languages, yet
- (ii) are not necessarily found in all languages, and
- (iii) differ in detail from language to language?

Although Mohanan raised these issues with regard to place assimilation, the same issue is also relevant to the case of NOCODA which was discussed in (32) through (35). Since Optimality Theory holds that constraints are violable in principle, it is expected that there will be languages which violate some constraint A and other languages that respect it.

To exemplify, let us return to the three outputs in (33). Both coda deletion in O_2 and vowel epenthesis in O_3 are strategies for avoiding codas. However, in a rule-based

approach, there is no formal way to express this relatedness. Deletion is expressed by a rule such as $C' \rightarrow \emptyset / _ \{C, \#\}$, while epenthesis requires a completely different rule, $\emptyset \rightarrow V / C' _$ (where C' represents an unsyllabifiable consonant). In Optimality Theory, on the other hand, the relation between these two processes can be captured through high ranking of NoCODA, as we have seen in (35). Optimality still requires mechanisms to ban epenthesis and deletion, DEP and MAX respectively. Thus, the main difference between the rule-based approach outlined above and a constraint-based approach like Optimality Theory is a difference in *focus*. While rule-based analyses focus on the ways that languages *resolve* constraints, i.e., on the rules themselves, Optimality Theory focuses on the underlying constraint, NoCODA, which may command epenthesis and/or deletion. Since in Optimality Theory, NoCODA has the same formal status as DEP and MAX, the ranking of the relevant faithfulness constraint (DEP or MAX) with respect to NoCODA will determine whether a language abides by NoCODA and, if it does, how it resolves NoCODA violations.

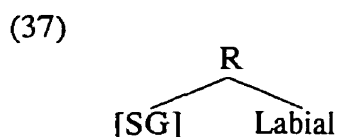
So far, we have looked only at the interaction between structural constraints like NoCODA and *segmental* deletion and epenthesis. However, since any constraint can be ranked with respect to any other, we should find that structural constraints also interact with *feature* parsing constraints like MAX [feature]. This type of interaction will be discussed at length in this thesis.

2.3.2 Optimality Theory and Feature Geometry

As mentioned earlier, most optimality-theoretic works which address alternations at the level of the segment make no direct reference to the hierarchical organization of features (Prince &

Smolensky 1993; Cole & Kisseberth 1994; Pulleyblank 1994, etc.). However, much of the earlier literature has shown the need to express constituency and dominance relationships, as discussed in Section 2.1.1. If we abandon feature geometry, we would need other ways of capturing these constituency and dominance effects. Expressing them through constraints is one option.¹⁷ However, since in Optimality Theory, constraints are rankable and violable in principle, we would expect constituency and dominance relations to be observed in a relative rather than absolute manner across languages. While some such relations may indeed vary across languages and thereby be best expressed through constraints, to abandon feature geometry altogether would surely create a problem of overgeneration of unattested grammars.¹⁸

I propose that feature geometry is encoded in GEN and that GEN only produces candidates which are licit in terms of the dependencies encoded in the geometry. Therefore, relations like those in (23) are respected by all outputs. If it is not accorded this formal status, the incorporation of feature geometry into the optimality-theoretic framework would not prevent a configuration such as that in (37) from being produced by GEN as a candidate.



Notice that the structure in (37) is not possible under either of the geometries in (2) or (3), nor under any version of feature geometry which has been proposed. The problems with (37) are: (a) [spread glottis] is directly dominated by Root, not by the Laryngeal node, and (b) Labial is not dominated by the Place node. Since (37) would never be selected as

optimal, allowing such a structure to be generated as a candidate would constitute a case of overgeneration.

2.4 The Role of [voice] in the Sonority Hierarchy

In the final section of this chapter, I will discuss the consequences that the feature geometry proposed here has for the interpretation of relative sonority. Since I have proposed that [voice] is doubly-dependent on both sonority nodes—SV and Laryngeal, it should also be relevant in determining the sonority value of segments. I will demonstrate that this is indeed the case.

The notion that sounds can be ranked in terms of sonority goes back to Whitney (1874). Since then, sonority hierarchies have been proposed to account for syllabification patterns across languages (Sievers 1885; Jespersen 1904; Saussure 1916; Hooper 1976; Greenberg 1978; Selkirk 1982, 1984, and others). Jespersen's (1904) version of the sonority hierarchy for consonants is given in (38).

(38) Jespersen's (1904) Sonority Hierarchy—Consonants

voiceless stops / voiceless fricatives
voiced stops
voiced fricatives
nasals / laterals
voiced r-sounds

A sonority scale like that in (38) may govern the phonotactics of adjacent speech sounds. Within a syllable, for example, the most sonorous segment is considered to be the *peak* and

sonority must decrease toward the syllable edges. Although sonority scales are argued to be universal, some segment classes may be collapsed in some languages. For example, all sonorant consonants may be collapsed as one class, and all obstruents may be collapsed as another.

Notice that in Jespersen's version of the sonority scale, voicing is relevant in defining sonority. More recently, however, it has been argued that voicing is not relevant for the determination of sonority (Clements 1990; Cho 1991, cf. Zec 1988). Zec (1988), for example, argues that placing both [continuant] and [voice] in the universal sonority scale yields conflicts across languages. If a language chooses only continuancy to be relevant for sonority and another language chooses only voicing, then /d/ would be less sonorous than /s/ in the former language, but the reverse would hold in the latter language. Zec therefore suggests that [voice] and other features such as [continuant], [coronal] and [anterior] are features which are added to the class of features which determine sonority on a language-particular basis.¹⁹

While I do not address the status of continuancy vis-à-vis the sonority hierarchy, I will demonstrate that there exist several cases which support the position that voicing is relevant for sonority. Greenberg (1978) observes that many languages which allow onset or coda obstruent clusters require the clusters to agree in voicing. However, if they do not agree in voicing, the sequence is always voiceless-voiced in onset and voiced-voiceless in coda. Greenberg's observation can be explained if voiced obstruents are more sonorous than voiceless ones.

Even in English, where voicing is argued not to play a role in determining sonority, there exist some distributional facts that suggest that voiced obstruents are more sonorous than voiceless ones. Across languages, it has been claimed that in hetero-syllabic consonant clusters, codas tend to be at least equal to or more sonorous than the following onset

(Murray & Vennemann 1983; Kaye, Lowenstamm & Vergnaud 1990, among others). In English, non-derived hetero-syllabic clusters and hetero-syllabic clusters derived by Level I morphology belong to the following classes: voiceless-voiceless, voiced-voiced and voiced-voiceless; the latter class is exemplified in (39). Hetero-syllabic voiceless-voiced obstruent clusters are found only in proper nouns (Lamontagne 1993); see (40).

(39) English: Voiced-Voiceless Obstruent Clusters
(Lamontagne 1993:313)

substitute	absurd	obscure	absorb
substance	absent	obsession	adscript
subscribe	abscess	obstruct	
subserve	abscind		
subsidence	absolute		
subsist	abstain		
subsume			

(40) English: Voiceless-Voiced Obstruent Clusters
(Lamontagne 1993:306)

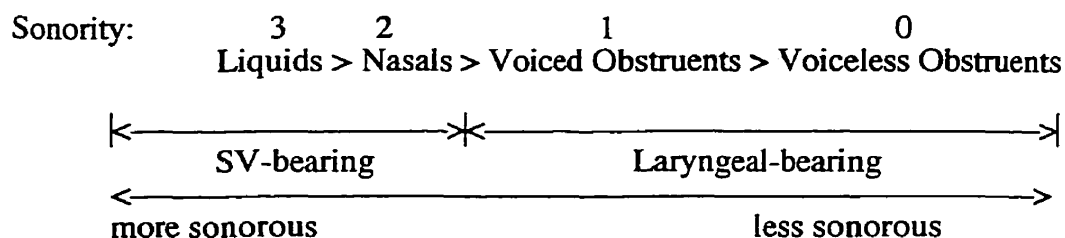
Afghan	Ashburn	Ashboro	Batesburg
Bronxville	Fitzgerald	Hepburn	Karlsbad
Lewisville	Macbeth	Updike	Oakdale
Pittsburg	Rathbone	Wolfgang	

Assuming that there is something special about names,²⁰ the English facts above are consistent with Greenberg's (1978) observation mentioned earlier. Given the tendency that a coda must be equal to or more sonorous than the following onset, the English facts above support the view that voiced obstruents are more sonorous than voiceless ones.

There are other languages as well where voicing seems to play a role in determining sonority, e.g., Irish (Carnie 1994), Attic Greek (Steriade 1982).²¹ Thus, although there is

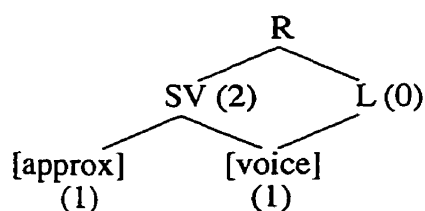
some disagreement over the inclusion of [voice] in the sonority scale, the evidence discussed thus far favors the view that voicing plays a role in determining the sonority value of segments. Furthermore, there are no languages where voiceless obstruents are more sonorous than voiced ones. In light of this, I provide the sonority scale for consonants in (41) below.²²

(41) Sonority Scale—Consonants



Although the scale in (41) expresses relative sonority in terms of values assigned to various types of segments, this is merely a notational device, as segments are not phonological primitives. It is reasonable to hypothesize that constraints can only make reference to phonological primitives, in this case to features. The values on the scale in (41) may be computed from values assigned to individual features and nodes as in (42).²³

(42) Sonority by Structure²⁴



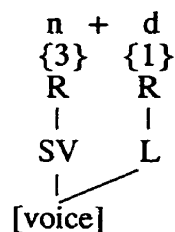
In the structure above, SV adds a value of 2. Dependents of SV each add a value of 1. Since Laryngeal is the node which defines obstruency, it contributes nothing to the sonority value of a segment.²⁵ Sonority values for each class of consonants as per (42) are provided in (43).

(43)

a. Liquids	b. Nasals	c. Voiced Obstruents	d. Voiceless Obstruents
R SV [approx]	R SV	R L [vce]	R L
Sonority Value 3	2	1	0

Explanations for the values assigned in (43) are as follows. First, recall from Section 2.1.1.1 that a bare SV node is interpreted as a nasal. Therefore, the feature [nasal] is not projected in the unmarked case; consequently, nasals have a lower sonority value than liquids.²⁶ Second, as discussed earlier, [voice] is dependent on both Laryngeal and SV. As a result, voiced obstruents are more sonorous than voiceless ones. Third, although [voice] is specified for sonorants underlyingly, this feature cannot be licensed by SV, as will be discussed in Section 4.2.2. Since the sonority values will be translated into a constraint on *outputs*, [voice] on SV will not normally feature into the sonority value calculation. It will, however, when [voice] for sonorants can be parasitically licensed by a following Laryngeal node. (44) below illustrates an instance of parasitic licensing.

(44) Parasitic Licensing of [voice]



In the parasitic licensing configuration in (44), the nasal which bears [voice] has a sonority value of 3 rather than 2 since the feature [voice] adds a value of 1. Apart from this type of configuration, however, sonorants do not bear [voice] for reasons that will be discussed in Section 4.2.2. As a result, nasals and approximants normally have sonority values of 2 and 3, respectively.

As mentioned earlier, hetero-syllabic consonant clusters are restricted in terms of their sonority relations. The restriction which requires codas to be equal to or more sonorous than the following onset has been proposed as the Syllable Contact Law by Murray & Vennemann (1983).²⁷ Based on the scale in (41), a set of sonority sequencing constraints between hetero-syllabic coda-onset clusters is formulated in (45) below.

(45) Syllable Contact Constraint (SONORITY)

In a hetero-syllabic consonant cluster, $C_1]_{\sigma\sigma}[C_2$, the sonority value of C_2 ($\{C_2\}$) minus the sonority value of C_1 ($\{C_1\}$) cannot be more than 0.

The constraint in (45) prohibits an onset from being more sonorous than the preceding coda. When the sonority value of the onset ($\{C_2\}$) is greater than the sonority value of the preceding coda ($\{C_1\}$), the value of $\{C_2\}-\{C_1\}$ is more than zero. Such a sequence violates SONORITY. I propose that the constraint in (45) is interpreted gradiently. For example, in the sequence “..t]_{σo}[n..” (where $t=0$ and $n=2$), $\{C_2\}-\{C_1\}$ is 2; the result is two violations of SONORITY. In the sequence “..d]_{σo}[n..” (where $d=1$ and $n=2$), $\{C_2\}-\{C_1\}$ is 1; therefore, one violation of SONORITY is incurred. This constraint has nothing to say about the relative well-formedness of clusters such as “..n]_{σo}[t..” vs. “..n]_{σo}[d..” because, for both, C_2-C_1 is less than zero.

The characterization of the SONORITY constraint proposed here is not equivalent to that proposed by Rice (1992). Since I assume that SV and Laryngeal are Sonority nodes and that [voice] is dependent on both of these nodes, both SV and Laryngeal contribute to the determination of relative sonority. This differs, both theoretically and empirically, from the structural sonority scale proposed by Rice where sonority is determined by the amount of structure under the SV node only.

In English, most hetero-syllabic clusters respect SONORITY in (45). Since voiceless + voiced obstruent sequences violate SONORITY by one, such sequences cannot be realized in English (but cf. (40)). However, violations are tolerated when the demands of higher ranked constraints must be met. For example, English has a /t + l/ hetero-syllabic cluster which violates SONORITY by three; e.g., ‘atlas’ [æt.ləs], *[æ.tləs]. These violations of SONORITY are forced because syllabifying both /t/ and /l/ into the onset would induce a violation of an even higher ranked constraint, one which prohibits onset clusters that contain consonants with identical place.

2.5 Summary

In this chapter, I have proposed that subsegmental structure must be part of an Optimality-theoretic grammar; more specifically, that feature geometry must be encoded into GEN. The combination of feature geometry and constraints yields a theory which is more constrained than one which relies on constraints alone. Because of this assumption, candidates which GEN produces are significantly limited, since GEN only produces structures which are consistent with the geometry.

With regard to subsegmental structure, I have adopted several hypotheses: 1) SV and Laryngeal are Sonority nodes, where SV defines sonorancy while Laryngeal defines obstruency; 2) the feature [voice] is dominated by both Sonority nodes.

In addition to subsegmental structure, I have also proposed the need for reference to higher prosodic structure and relations that hold between syllable positions. By adopting the notion of “head” from Government Phonology, I will account for why licensing relations are not bidirectional.

In the following chapter, I will introduce two coda constraints; a constraint on coda Laryngeal, and a constraint on coda Place. I will then demonstrate how the proposed feature geometry, together with these constraints, accounts for various coda phenomena across languages.

—NOTES—

¹ See Steriade (1995) for an overview of this issue.

² In this geometry, the organization of place features is omitted since the focus here is on sonority and laryngeal features. Throughout this thesis, only relevant parts of structures are given and irrelevant parts are omitted.

³ However, there are other types of nasal harmony, what Piggott (1992a, b) calls Type A harmony (e.g., Urhobo, Warao, Malay). In this type of nasal harmony, obstruents (and in some languages, liquids or glides also) serve as blockers of the harmony. To account for these languages, Piggott (1992a, b) has proposed that [nasal] can variably be a dependent of SV or of SP (Soft Palate). For detailed discussion, see Piggott (1992a, b).

⁴ This representation does not reflect Rice and Avery's (1991) analysis. They argue that only features and not nodes can spread. According to their analysis, the process in (9) involves two steps: copying of the SV node and spreading of the feature [nasal] or [lateral].

Since I assume that nasal is the default interpretation of SV following Rice (1993), the nasal in (9) does not bear the feature [nasal] (see below in text).

⁵ Coda constraints provide further support for nasals lacking SV dependents. See Chapter 3.

⁶ Although some works use different terms to refer to these features (e.g., [aspirated] for [spread glottis]), the structure in (16) reflects the standard view of Laryngeal organization. A couple of other geometries have also been proposed. For example, Ladefoged (1989) has argued for a more complex organization where Laryngeal dominates Voice, Glottal Aperture, Aspiration and Pitch, and each of these features in turn dominates sub-features such as [voice], [closed], [creaky], etc. For details, see Ladefoged (1989).

⁷ This problem is recognized by Rice (1993), Itô, Mester and Padgett (1995) and others. Their proposals are addressed in Section 5.1.1.

⁸ The idea of dual-dependency (or variable dependency) of features has been suggested for other features as well. For example, place features are proposed to be doubly-dependent on C-Place and V-Place by Clements and Hume (1995). Recall as well from endnote 3 above that Piggott (1992) proposed that [nasal] is doubly-dependent on SP and SV.

⁹ In an approach which does not assume SV, sonorants and voiced obstruents can form a natural class—as segments which bear the feature [voice]. However, such an approach is argued to be insufficient in accounting for cases where [voice] for sonorants is phonologically inert (e.g. *Rendaku*). Within the framework of Underspecification Theory, such cases have been handled by redundancy rule ordering (e.g., Kiparsky 1985; Itô & Mester 1986). [voice] for sonorants is not present underlyingly and a redundancy rule, [sonorant]→[voice], applies at a later stage in the phonology. However, there are empirical problems with this type of approach as discussed in Itô, Mester & Padgett (1995). See further Section 4.2.

¹⁰ In (25), place assimilation is omitted. The square around Laryngeal indicates that L is inserted.

¹¹ For a comprehensive comparison between Moraic Theory and Onset-Rhyme theory (which is shown in (26)), see Hayes (1989), Blevins (1995) and Broselow (1995).

¹² In standard moraic theory, the moraic tier also replaces the timing tier (X or C/V) which mediates between the Root node tier and the terminal syllable structure constituents.

¹³ These two families of constraints are correspondence-based versions of PARSE and FILL respectively in the original work by Prince and Smolensky (1993).

¹⁴ Although I accept the Optimality Theory premise that rules and derivation have no formal status in the theory, I use the words ‘delete’, ‘insert’ and ‘spread’ for convenience.

¹⁵ This does not mean to suggest that there are only two types of constraints. There are other types of constraints such as “alignment” (McCarthy & Prince 1993). Alignment requires the edge of some prosodic or morphological category to be “aligned” with the edge of some other category.

¹⁶ In the tableaux, constraint violations are marked by * and the optimal candidate, that which best-satisfies the constraint-ranking, is indicated by ✱. Dotted lines between constraints indicate that the ranking is indeterminable.

¹⁷ Padgett (1995), for example, proposes “feature class” theory within the framework of Optimality Theory. Under this proposal, constituency is expressed through reference to feature classes. For example, the laryngeal features [SG], [CG] and [voice] constitute the “Laryngeal” class. In contrast to feature geometry, the notion of “class” is not encoded in terms of hierarchical structure. Regarding subsegmental structure, he reverts to the “bottle-brush” model of early autosegmental phonology where almost all features and nodes (except for [anterior] and [distributed] which are dominated by Coronal) are directly dominated by Root. Although Feature Class Theory can capture the tendency that some features often pattern together in phonological processes, it is very unrestricted in terms of the number of possible representations it allows. Therefore, it also suffers from the overgeneration problem mentioned below in the text.

¹⁸ Feature Geometry is by no means perfect, given that many different models have been proposed which continue to be revised. However, certain general properties consistently re-emerge (e.g., reference to the Place and Laryngeal nodes). This suggests that these properties are absolute, and not violable. This thesis demonstrates that combining feature geometry and constraints is more restricted in its predictions than encoding constituency in terms of violable constraints alone.

¹⁹ Cho (1991), on the other hand, argues that [voice] is not relevant for sonority at all. Cho focuses on the observation that voiced-voiceless obstruent sequences in onset have not been attested in any language while voiceless-voiced sequences are attested in some languages (Coeur d'Aléne, Palaychi Karen). Cho explains the lack of the former type of sequence by a constraint which bans a voiced obstruent before a non-sonorant consonant within an onset (Universal Devoicing) (Cho 1991:72). Although his constraint accounts for the absence of voiced-voiceless obstruent sequences in onset, the constraint is merely description of the distributional facts and is not independently motivated. Moreover, it will not account for the heterosyllabic tendencies discussed below in the text.

²⁰ For many speakers, the words in (40) exhibit compound stress, so these examples may not even be exceptions.

²¹ In Attic Greek, onset clusters are either voiceless stops + [n, l, r] or voiced stops + [l, r]. Clusters like voiced stop + [n] are not licit. From this fact, Steriade (1982) suggests that voiced obstruents are more sonorous than voiceless ones. Other supporting evidence from Ukrainian will be provided in Chapter 4.

²² I do not include glides in the hierarchy since glides are vowel-like segments and their sonority status is not well understood. The sonority status of glides is left for further research.

²³ Earlier proposals which encode the sonority hierarchy in terms of features are made by Selkirk (1984) and Clements (1990). For other hierarchical accounts of relative sonority, see Harris (1990) and Rice (1992).

²⁴ As discussed in Section 2.1.1.1, the feature [nasal] is assumed to be absent in the unmarked case. When [nasal] is present, it will also contribute to the calculation of sonority ([nasal] adds a value of 1).

²⁵ Government Phonology (e.g., Kaye, Lowenstamm & Vergnaud 1985) expresses segmental strength in terms of ‘charm’ which is similar to the sonority values proposed here.

²⁶ In the marked case where nasals are specified as [nasal], nasals and liquids would be equal in sonority.

²⁷ Murray & Vennemann (1983:520) formulate the Syllable Contact Law in terms of consonant “strength”, which is roughly equivalent to sonority:

The Syllable Contact Law: The preference for a syllabic structure AB , where A and B are marginal segments and a and b are the Consonantal Strength values of A and B respectively, increases with the value of b minus a .

3 . VOICING AND CODA CONSTRAINTS

In the previous chapter, I introduced the constraint NOCODA which militates against codas. The original formulation of this constraint in Prince & Smolensky (1993) accounts for languages which do not allow codas at all. However, NOCODA by itself is far from satisfactory in explaining the fact that languages may place different types of restrictions on codas. This chapter deals with several constraints on codas, focusing mainly on Laryngeal specifications. First, I will explore the cross-linguistic restrictions on coda Laryngeals and will show how these phenomena can be captured within the framework of Optimality Theory. Secondly, I will show that sonority requirements on codas which some languages exhibit can be accounted for by a coda Laryngeal constraint, NOCODA: LARYNGEAL. Third, I will demonstrate that three types of languages— 1) languages which allow only sonorants in coda, 2) languages which allow only nasals in coda, and 3) languages which allow only glottal stops in coda— can be captured in a unified way through NOCODA: LARYNGEAL. Lastly, I will discuss laryngeal assimilation cases which I argue are also a consequence of NOCODA: LARYNGEAL.

3.1 Coda Constraints

To account for the absence of codas in some languages, we have shown how the different rankings of Faithfulness constraints (MAX and DEP) and the NoCODA constraint select different outputs for input post vocalic consonants in Section 2.1.2. NoCODA bans the appearance of any segment in coda position. As mentioned in the introduction, however, there are also languages where only certain consonants are found in coda position. I will refer to all such restrictions as *Coda Constraints*. The strongest among the Coda Constraints is NoCODA. In this section, I will introduce two more Coda Constraints: NoCODA: LARYNGEAL (*CODALAR), and NoCODA: PLACE (*CODAPL). These two constraints prohibit codas from bearing a Laryngeal node and a Place node, respectively. Together, these constraints account for common phenomena such as coda devoicing, coda deglottalization, and place assimilation.

3.1.1 NoCoda: Laryngeal

The neutralization of laryngeal features in coda has recently been of issue (Cho 1990a, Lombardi 1991, 1995b, among others). Voicing distinctions are neutralized in coda in many languages (e.g. German, Catalan, Kirghiz, etc.). In addition, other laryngeal features (e.g., [spread glottis] ([SG]) and [constricted glottis] ([CG])) are also restricted in coda position. For example, Maidu, which allows voiced, glottalized and plain (voiceless unaspirated) obstruents in onset, only allows the plain obstruents in coda (Shipley 1956, 1963). Examples from German and Maidu are provided in (1) and (2) respectively.

(1) German (Mascarô 1987)

a.	run[d]e	'round (pl)'	We[g]e	'way (Dat)'
b.	run[t]	'round (sg)'	We[k]	'way (Nom)'
c.	Run[tg]ang	'round'	we[kz]am	'transitable'
d.	Run[ts]äule	'cylinder'	We[k s]pur	'trace'

(2) Maidu (Shipley 1963)

a. /lût'/	'real, the real one, completely, extremely, the ...est'
tibílut'i	'the smallest'
nenólùt	'eldest'
tetélùt	'enormous'
b. /dýk'/	'only, alone, just ... and no more'
c'ak'ámdyk'ý	'pitch: "just pitch and nothing more"'
?adýk	'furthermore, just so'
?ypék'andýk	'all, every single, every last one'

As the German data in (1) show, /d, g/ in (1a) become voiceless /t, k/ when they are syllabified as codas as in (1b-d). The data in (2) demonstrate that Maidu neutralizes [CG] contrasts; when a stem-final glottalized obstruent is syllabified as a coda, it loses glottalization (/lût'/ → [nenólùt] 'eldest'). Shipley notes, in addition, that voiced obstruents are not permitted in coda, although he provides no morphological alternations which show this.

In addition to Maidu, Kiowa is another language which neutralizes multiple laryngeal features in coda. See the Kiowa examples below.

(3) Kiowa (Watkins 1984)

Kiowa Consonant Inventory

p'	t'	k'
p ^h	t ^h	k ^h
p	t	k
b	d	g
	c'	
	c	
	s	h
	z	
m	n	
	l	
	y	

Kiowa Onsets

a.	p'í	'female's sister'	(p. 7)
	p ^h í	'fire; hill; heavy'	(p. 7)
	pó	'eat'	(p. 7)
	bó	'bring'	(p. 7)

Kiowa Codas

b.	kól	'bison cow'	(p. 29)
	cégùn	'dog'	(p. 21)
	cát	'entrance, doorway'	(p. 20)
		*cád *cát' *cát ^h	
	t'áp	'deer'	(p. 7)
		*t'áb *t'ap' *t'ap ^h	

As the consonant inventory in (3) shows, [voice], [SG], and [CG] are all contrastive in Kiowa. (3a) reveals that all laryngeal contrasts may appear in onset. However, possible coda consonants are either sonorants or plain voiceless obstruents as in (3b). Among the obstruents, the voiced, aspirated, and glottalized segments never appear in coda.

In the languages discussed above, all laryngeal contrasts that a given language permits are neutralized in coda. There are also languages which neutralize only a subset of laryngeal contrasts in this position. According to Fleming & Dennis (1977), Tol, which has [SG] and [CG] distinctions, allows only [CG] in coda. See the description in (4).

(4) Tol (Fleming & Dennis 1977:122)

“Initially in a syllable, there is a three-way contrast of plain, aspirated, and glottalized stops. Syllable-finally only a two way contrast has been found, glottalized and nonglottalized.”

Similarly, in Gujarati, where [voice] and [SG] are contrastive in onset, only [voice] is allowed in coda.

We have seen that different laryngeal contrasts are neutralized in coda, depending on the language. Furthermore, some languages (e.g., English) do not exhibit any neutralization of laryngeal features in coda. Optimality Theory’s provision for constraint violability permits us to provide a parsimonious solution to this observation by collapsing these separate constraints into one: NOCODA: LARYNGEAL as in (5).

(5) NOCODA: LARYNGEAL (*CODALAR)¹

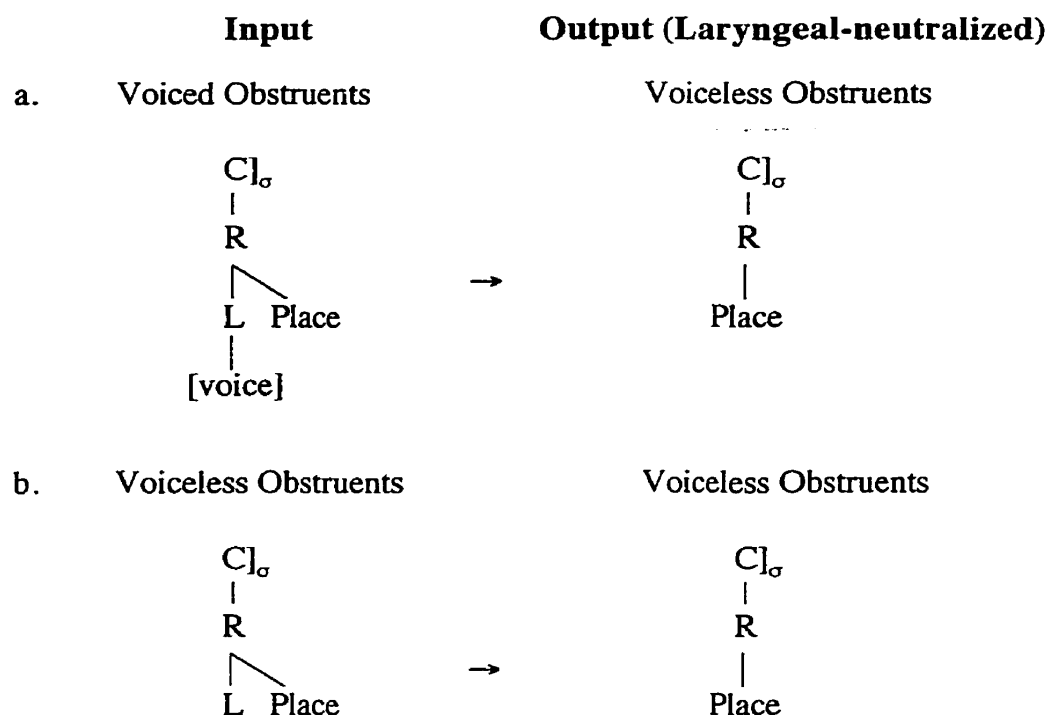
Codas cannot license a Laryngeal node.

$$\begin{array}{c} *C]_{\sigma} \\ | \\ L \end{array}$$

The constraint in (5) limits the material that can be *licensed* in coda. Since all nodes and features must be licensed in order to be phonetically realized (linked) or parsed (cf. Itô’s (1986) Prosodic Licensing), a coda consonant which exhaustively bears a Laryngeal node

will violate *CODALAR since the Laryngeal node in such a segment must be licensed by the coda. With *CODALAR, voicing neutralization in languages like German is the result of unparsing (delinking) the coda Laryngeal node to be faithful to *CODALAR. See the representations in (6) below where $C]_{\sigma}$ represents the coda position.

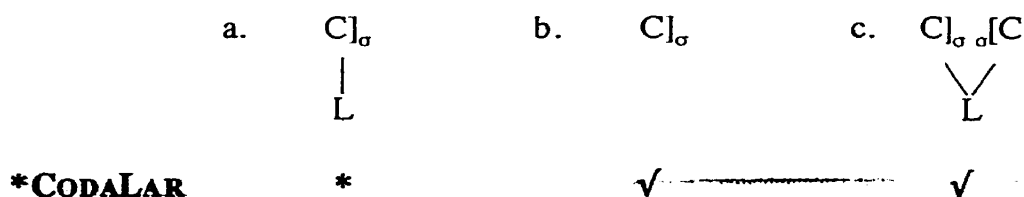
(6) Laryngeal Neutralization



As (6a) illustrates, voiced obstruents become voiceless in coda by losing Laryngeal. Although it does not bear any features below, Laryngeal is specified for a voiceless obstruent in the input,² but I argue that Laryngeal is missing from output voiceless coda consonants in languages like German, Kiowa and Maidu. Although voiceless obstruents do not appear to be changed in the output, I contend that all Laryngeal contrasts are lost in coda in these languages.

However, there is one configuration under which codas can bear Laryngeal without violating *CODALAR. Consider the configurations in (7), where $C]_{\sigma}$ represents the coda position.

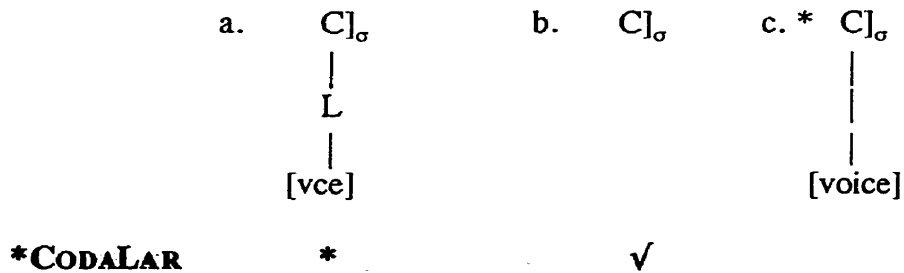
(7) Licensing vs. Parasitic-Licensing



The structure in (7a) violates *CODALAR since the Laryngeal node is exhaustively linked in coda. However, *CODALAR is not violated by either of the structures in (7b) and (7c). In (7b), there is no Laryngeal node present, so *CODALAR is clearly satisfied. In (7c), the Laryngeal node belongs to both the coda and the onset. In such a configuration, only one of the two positions needs to license the Laryngeal specification. Since the onset can serve as the licenser of Laryngeal, *CODALAR is not violated. This is an instance of *parasitic licensing* (cf. indirect licensing in Steriade (1995)). We will return to parasitic licensing in Chapter 4.

Since *CODALAR prohibits a Laryngeal node from being licensed in coda, a Laryngeal node which is not parasitically licensed is lost in languages which abide by this constraint. When [voice] is a dependent of this node, it is also lost (unless it migrates to be licensed by another position). Compare (8a) and (8b).

(8) Loss of [voice] due to *CODALAR



For [voice] to be exhaustively parsed in coda, it must be dominated by Laryngeal node as in (a). However, (a) violates *CODALAR. To be faithful to *CODALAR, [voice] must be lost together with Laryngeal. Since feature geometry, which is assumed to be in GEN, requires a Laryngeal node or SV node in order to parse [voice], GEN will not produce a candidate where [voice] is immediately dominated by the Root node, (8c).

I will now show how different types of laryngeal feature neutralization can be accounted for in terms of the constraint *CODALAR. The constraints which interact with *CODALAR belong to the MAX family of constraints.

(9) MAX Constraints for laryngeal features

MAX [voice] (MAXVCE)

A feature [voice] in the input has a correspondent in the output.

MAX [spread glottis] (MAX[SG])

A feature [spread glottis] in the input has a correspondent in the output.

MAX [constricted glottis] (MAX[CG])

A feature [constricted glottis] in the input has a correspondent in the output.

These three constraints prohibit any laryngeal feature which is specified in the input from being deleted in the output. If all of the MAX constraints for laryngeal features outrank *CODALAR, no laryngeal features will be lost in coda; see the tableau in (10a). Candidate 1, for which the laryngeal feature α is preserved, will be selected as optimal. If *CODALAR outranks all of the MAX constraints in (9), all laryngeal features will be neutralized in coda. In (10b), candidate 3 is selected since it is the only candidate which satisfies the highest ranked constraint, *CODALAR. Thus, we have seen that the interaction of these MAX constraints with *CODALAR yields different optimal outputs. Notice that no matter how we rerank the constraints in (10), candidate 2 will never be selected as optimal.

(10) Constraint reranking: MAX, *CODALAR

Input = $\begin{matrix} C]_{\sigma} \\ | \\ R \\ | \\ L \\ | \\ \alpha \end{matrix}$

(where α represents any laryngeal feature)

a. MAXCG, MAXSG, MAXVCE » *CODALAR

	MAXCG, MAXSG, MAXVCE	*CODALAR
1. $\begin{matrix} R]_{\sigma} \\ \\ L \\ \\ \alpha \end{matrix}$		*
2. $\begin{matrix} R]_{\sigma} \\ \\ L \end{matrix}$	*	*
3. $\begin{matrix} R]_{\sigma} \end{matrix}$	*	

b. *CODALAR » MAXCG, MAXSG, MAXVCE

	*CODALAR	MAXCG, MAXSG, MAXVCE
1. $\begin{matrix} R]_{\sigma} \\ \\ L \\ \\ \alpha \end{matrix}$	*	
2. $\begin{matrix} R]_{\sigma} \\ \\ L \end{matrix}$	*	*
3. $\begin{matrix} R]_{\sigma} \end{matrix}$		*

If only a subset of the MAX constraints for laryngeal features ranks lower than *CODALAR, only a subset of laryngeal features will be neutralized in coda. For example, if MAXCG ranks higher than *CODALAR while the other MAX constraints for laryngeal features rank lower than *CODALAR, only [CG] in coda will be retained and other laryngeal features, [voice] and [SG], will be lost; see (11).

(11) Constraint ranking: MAXCG » *CODALAR » MAXSG, MAXVCE

inputs: a. $C]_{\sigma}$
 R
 |
 L
 |
 [vce]
 b. $C]_{\sigma}$
 R
 |
 L
 |
 [SG]
 c. $C]_{\sigma}$
 R
 |
 L
 |
 [CG]

		MAXCG	*CODALAR	MAXSG	MAXVCE
a.	1. $R]_{\sigma}$ L [vce]		*		
	2. $R]_{\sigma}$				*
b.	1. $R]_{\sigma}$ L [SG]		*		
	2. $R]_{\sigma}$			*	
c.	1. $R]_{\sigma}$ L [CG]		*		
	2. $R]_{\sigma}$	*			

As the tableaux in (11) show, only in (11c) will the candidate where the Laryngeal contrast is not neutralized be selected. In the other tableaux, the candidates where Laryngeal is maintained, (11a-1) and (11b-1), violate *CODALAR while satisfying the lower ranked constraints, MAXVCE and MAXSG. Therefore, the candidates where Laryngeal is lost will not be selected for these tableaux.

As (11) demonstrates, this constraint-based approach to coda neutralization accounts for the existence of languages where some laryngeal features are lost while others are retained. As discussed earlier, in Tol, which has [SG] and [CG] distinctions, only [SG] is neutralized while [CG] persists. In addition, in Gujarati, where both [voice] and [SG] are contrastive in onset position, only [SG] contrasts are lost syllable-finally.

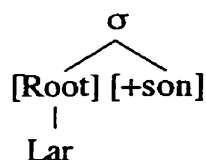
3.1.1.1

Voicing Neutralization

In the previous section, I proposed that a single constraint is responsible for coda Laryngeal neutralization. Its interaction with independently needed Max constraints yields different types of languages: languages where all laryngeal contrasts are lost in coda, languages where all laryngeal contrasts persist, and languages where only a subset of laryngeal contrasts is lost.

In earlier work, Lombardi (1991, 1995b) takes a different approach. To account for the frequent occurrence of neutralization of [voice] as well as other laryngeal features, she proposes the Laryngeal Constraint in (12) below.

(12) Laryngeal Constraint (Lombardi 1991, 1995b)



This constraint holds that a Laryngeal node can only be licensed in pre-sonorant position within a syllable. Therefore, the presence of Laryngeal in coda, which does not conform to the configuration in (12), violates this constraint.

The Laryngeal Constraint in (12) is a bipositional constraint which refers to a specific sequence (Laryngeal + sonorant) in a specific environment (within a syllable). With this formulation, Lombardi is able to capture two types of restrictions: (a) coda Laryngeal neutralization, and (b) voicing agreement in obstruent clusters within a syllable. The latter is observed in languages like Polish. The bipositional constraint formulation in (12) therefore seems to have an advantage over the NOCODA: LARYNGEAL constraint introduced in Section 3.1.1.

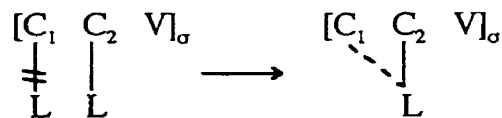
Let us turn to the Polish data cited by Lombardi which show that obstruent clusters are either uniformly voiced or uniformly voiceless.

(13) Polish (Lombardi 1995a)

[gd]y	‘when’	o[dg]rodzić	‘separate’
[db]ać	‘take care’	gwia[zd]a	‘star’
[bzd]ura	‘nonsense’	o[dvz]ajemić	‘reciprocate’
[dždž]ownica	‘earthworm’		
[pt]ak	‘bird’	ne[ptk]a	‘small frog’
ró[zg]a	‘rod’	pa[štš]a	‘gorge’
[pštš]oła	‘bee’	gwia[stk]a	‘star, dim’
[pst]ry	‘gaudy’	o[tst]raszyć	‘scare’

Under Lombardi’s bipositional constraint, voicing agreement in Polish clusters can be accounted for as follows. Since only an obstruent in pre-sonorant position can license a Laryngeal node, obstruents which are in other positions lose their Laryngeal nodes and, in turn, receive the Laryngeal specification of the following obstruent through spreading. This is illustrated in (14).

(14) Laryngeal Sharing



With regard to the realization of coda obstruents, both the bipositional constraint in (12) and NoCODA: LARYNGEAL account for the devoicing phenomenon discussed in Section 3.1.1. While NoCODA: LARYNGEAL in (5) cannot be extended to the onset voicing agreement

in (13), I nevertheless argue that it should be adopted. As Lombardi (1991) herself points out, if we allow the constraint to be formulated as in (12), then changing the primitives will yield unattested constraints such as the following: ‘Laryngeal can be licensed only before obstruents’, ‘Laryngeal can be licensed only after sonorants’, etc. To counter this argument, Lombardi claims that “the constraint that I propose is not composed of such options. This constraint is trying to capture an insight that is not composed of our current theoretical primitives, which combine with options about directionality and edges. Its form arises from deeper principles that are not subject to this variation. The principles have something to do with [voice] in obstruents only being able to appear before a more sonorous segment (vowel or sonorant).... [T]he transition from consonant to vowel (or to other sonorous segment) has special properties. There is something important about this particular transition, in a way that does not admit of using the left/right distinction to cross-classify it, the way the left/right distinction can cross-classify word edges or sonority (towards the nucleus/away from the nucleus). The transition from vowel to consonant is not comparable; directionality is not an option, because it is the particular direction that is crucial and that has special properties.” (pp. 79-80). The “special pleading” in this long quotation cannot be arrived at from the formulation in (12); it can only be determined from the explication of (12). From Lombardi’s statement, it is not clear what the ‘deeper principles’ are; thus, there is no way for us to predict what other constraints are possible as extensions of the constraint she proposes. Also, since the constraint in (12) is bipositional, without incorporating a notion of *head*, the ‘particular transition’ that the constraint is supposed to capture is not expressed in the formulation.

The formulation of NoCODA: LARYNGEAL in (5), on the other hand, is consistent with the cross-linguistic generalization that languages tend to restrict what can appear in coda. Within some frameworks, for example *Government Phonology* (Kaye, Lowenstamm &

Vergnaud 1985, 1990, Kaye 1990, Harris 1990, 1994), phonological principles have been proposed which attempt to capture why codas are relatively restricted when compared with onsets. Since the coda is governed by both the nucleus and the following onset, it is dependent on other positions (Coda Licensing) (Kaye 1990). Together with NoCODA and NoCODA: PLACE (Itô 1986) (see Section 3.1.2), NoCODA: LARYNGEAL is one of the constraints which is consistent with the notion of government.

Empirically, the bipositional constraint appears to have an advantage in that it can account for voicing agreement in syllable-internal obstruent clusters. However, onset and coda obstruent clusters are rare outside of Slavic languages.³ Notice that these clusters violate the sonority requirement which states that sonority must rise in onset clusters and fall in coda clusters. Therefore, it is not clear whether or not the voicing agreement observed in this rare type of cluster should be treated in the same way as coda devoicing.⁴

3.1.2 NoCoda: Place

Similar to laryngeal features, place specifications in coda are also subject to restrictions. Some languages only allow coronals (which are often analyzed as placeless) or one half of a geminate to appear in coda (e.g. Finnish); others allow only the first half of a geminate, nasals which are homorganic with the following onset, or placeless nasals (e.g. Japanese). These crosslinguistic facts suggest that a coda constraint on Place is required. Parallel to *CODALAR, I adopt Itô's (1986) coda place constraint which is formulated as in (15) below.

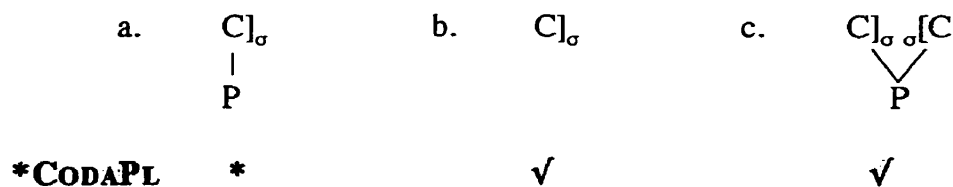
(15) NoCODA: PLACE (*CODAPL)

Codas cannot license a Place node.



This constraint prohibits codas from licensing a Place node. Therefore, the appearance of a singly-linked segment like [b] in coda violates *CODAPL. However, a homorganic segment whose Place node is also linked to the following onset circumvents the constraint. See the representations in (16).

(16)



In the representations above, (16a) violates *CODAPL since the coda licenses a Place node. On the other hand, (16b) which does not bear Place clearly satisfies *CODAPL. Contrary to appearances, *CODAPL is also satisfied in (16c) because Place is not licensed by the coda, but instead, by the following onset. Thus, homorganic clusters escape any violation of *CODAPL.

3.2 Coda Sonority Requirement

So far, we have seen that languages have coda restrictions on laryngeal and place features. However, restrictions imposed on codas are not limited to these features. Itô (1986) points out that there are some languages where only sonorants are allowed in coda. The Beijing dialect of Chinese⁵ (Blevins 1995) and Diola Fogy (Sapir 1969) are two such languages. To account for languages such as these, Itô (1986) formulates the coda condition shown in (17).

(17) Coda Sonority Constraint

$$\begin{array}{c} *C]_{\sigma} \\ | \\ [-son] \end{array}$$

This constraint prohibits any singly linked obstruent segment from occupying the coda position.

While Itô's coda sonority constraint in (17) can account for languages like Diola Fogy, it must be modified to capture properties of other languages because languages fall roughly into three types with regard to sonority restrictions: 1. languages which only allow sonorants (e.g., Beijing Chinese, Diola Fogy, Hausa), 2. languages which only allow nasals (e.g., Axininca Campa, Kiribatese, Japanese⁶), and 3. languages which only allow either /h/ or /ʔ/ and sonorants (e.g., Buginese, Macushi). I will demonstrate that the proposals presented so far can account for these three types of languages through a

combination of *CODALAR and other independently needed constraints.⁷ We will consider each case in turn.

3.2.1 Coda—Sonorant Only

NOCODA: LARYNGEAL prohibits Laryngeal from being licensed in coda. If the laryngeal node defines obstruency, it naturally follows that this constraint prohibits singly-linked obstruents in coda. I will demonstrate that an independent constraint which requires codas to be *sonorant* is not needed. As mentioned above, in some languages which abide by *CODALAR, codas are limited to sonorants. Diola Fogny is a language of this type; some examples are provided in (18).

(18) Diola Fogny (Sapir 1969)

- | | | |
|----|---------|--------------------|
| a. | ndaw | ‘or’ |
| b. | ekumbay | ‘the pig’ |
| c. | jɛnsu | ‘undershirt’ |
| d. | salte | ‘be dirty’ |
| e. | aɾɿ | ‘negative’ |
| f. | ijaut | ‘I did not come’ |
| g. | famb | ‘annoy’ |
| h. | kaŋg | ‘be furthest away’ |

In Diola Fogny, while word-final position may be occupied by consonants of any quality,⁸ word-internal codas are restricted to sonorants (/w, m, n, l, r/).

There are additionally some languages where coda obstruents become sonorants. Hausa (Hayes 1986) is one such language.

(19) Hausa (Hayes 1986:333)

a. sabroo	→	sawroo	'mosquito'
b. biyad	→	biyaṛ	'fine'
c. batagyee	→	batawyee	'twin'

As the examples in (19) show, when /b/ is in coda, it becomes a sonorant, [w]. Similarly, /d/ and /g/ become [ṛ] and [w] respectively. This process is called Klingenberg's Law.

(20) Klingenberg's Law (Hayes 1986)

$$[-\text{cont}] \rightarrow [+son] / \begin{array}{c} C|_s \\ \underline{\quad} \end{array}$$

In the framework developed here, Klingenberg's Law is expressed not only as loss of the Laryngeal node, but also as insertion of an SV node. SV insertion is motivated by a constraint which requires all segments to bear at least one sonority node, either Laryngeal or SV.⁹

(21) Sonority Node Requirement (SONNODE)

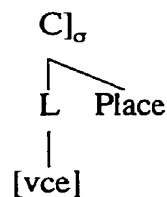
Every segment must be specified for sonority.
(Every segment must bear either a Laryngeal node or an SV node).

Importantly, however, the loss of the Laryngeal node does not always lead to the sonorantization of obstruents. For example, consider the German case in (1) above. Recall that in German, voiced obstruents become their voiceless counterparts when they occupy the coda position; they do not become sonorants. I suggested earlier that the surface representation for German coda obstruents does not contain a Laryngeal node, but these segments do bear a Place node. This hypothesis is consistent with the observation that in many German-type languages (e.g., Thai, Taiwanese), coda obstruents are unreleased. Thus, the surface representation of a devoiced obstruent does not contain either of the sonority nodes.

Representations for German/Thai-type languages and Diola Fogny/Hausa-type languages are in (22a) and (22b) respectively.

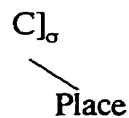
(22) *CODALAR

Input:

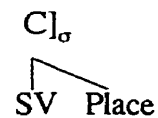


Outputs:

a. Voiceless Obstruent
(Unreleased)



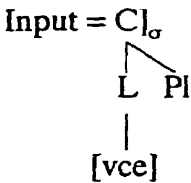
b. Sonorant



In languages which do not respect SONNODE, the result of satisfying *CODALAR will be the representation in (22a). This representation, which does not contain any sonority nodes, will be interpreted as a voiceless (unreleased) obstruent as in German and Thai. If a language respects SONNODE, the result of satisfying *CODALAR will be the representation in (22b)

where SV is inserted in place of Laryngeal as in Hausa.¹⁰ This output violates DEPSV, a member of the DEP family of constraints which, recall from (31) in Chapter 2, militates against the insertion of new features. The constraint-rankings for the two types of languages are given in (23).

(23) Constraints: *CODALAR, SONNODE, DEPSV



1. Diola-Fogny / Hausa-type:

	*CODA LAR	SON NODE	DEP SV
a. $C _{\sigma}$ / \ L Pl [vce]	*		
b. $C _{\sigma}$ / \ L Pl	*		
c. $C _{\sigma}$ \ Pl		*	
d. $C _{\sigma}$ / \ SV Pl			*

2. German / Thai-type:

	*CODA LAR	DEP SV	SON NODE
a. $C _{\sigma}$ / \ L Pl [vce]	*		
b. $C _{\sigma}$ / \ L Pl	*		
c. $C _{\sigma}$ \ Pl			*
d. $C _{\sigma}$ / \ SV Pl		*	

Since candidates (a) and (b) both violate the higher ranked constraint, *CODALAR, they lose out to the other candidates in both tableaux. In tableau 1, SONNODE is ranked higher than DEPSV. Therefore, (d) is selected as the optimal output. In tableau 2, on the other hand, DEPSV is ranked higher than SONNODE, so candidate (c), which respects DEPSV, is optimal.

Thus, the ranking between DEPSV and SONNODE is crucial in determining whether a language chooses voiceless obstruent or sonorant codas in order to be faithful to *CODALAR.

3.2.2 Coda—Nasal Only

If *CODALAR and SONNODE are both ranked higher than DEPSV in a language, the language should choose sonorants over obstruents in coda, as discussed earlier. However, many of these languages only allow nasals in coda,¹¹ e.g. Japanese, Kiribatese (Groves, Groves & Jacobs 1985), Axininca Campa (Payne 1981). Examples from Kiribatese and Axininca Campa are shown in (24) and (25), respectively.

(24) Kiribatese Codas (Groves, Groves & Jacobs 1985)

- | | | |
|----|---------|------------------------|
| a. | kaŋ | 'to eat' |
| b. | anti | 'ghost' |
| c. | naŋkiro | 'to be about to faint' |

(25) Axininca Campa Codas (Payne 1981)

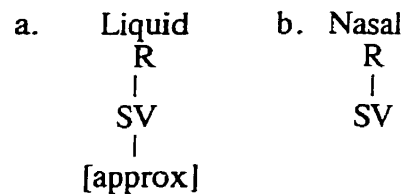
- | | | |
|----|------------------------|-----------------|
| a. | hanto | 'there' |
| b. | ts ^h amanto | 'woodpecker' |
| c. | impisiti | 'he will sweep' |
| d. | amimpori | 'white condor' |
| e. | noŋkimi | 'I will hear' |
| f. | iŋki | 'peanut' |

The fact that, among sonorant consonants, codas are sometimes restricted to nasals should make us suspicious about the definition of the coda restriction only in terms of sonority. If codas are required to bear SV, then why, among sonorants, are only nasals allowed in coda

in a number of languages? Since liquids are considered to be more sonorous than nasals, shouldn't they make better codas than nasals?

Recall from Chapter 2 that the difference between liquids and nasals is captured through the features which are specified under the SV node. Liquids bear the feature [approximant] under SV. Nasals, on the other hand, are represented as bare SV. Compare the two representations in (26).¹²

(26) SV structures for liquids and nasals



Thus, the difference between liquids and nasals can be attributed to a difference in the amount of structure under the SV node. To account for why nasals are preferred over liquids, I will propose the following constraint.

(27) No Dependency¹³
(*DEPEND)



An organizing node¹⁴ cannot bear a dependent feature.

*DEPEND in (27) prohibits any organizing node (SV, L, or Place) from bearing a dependent feature. It is a member of the No Complex Structure constraint family; another member

*COMPLEX will be introduced in Section 4.2.4.3. The No Complex Structure constraint family militates against segmental complexity; it is thus in the spirit of *STRUC proposed by Chomsky (1986, 1989) which bans unnecessary branching.¹⁵

*DEPEND receives support from the literature on underspecification. For example, it has been widely argued that coronal is the unmarked place of articulation and this is most often expressed through coronal underspecification (see e.g., Paradis & Prunet 1991b). If coronal is represented as a bare Place node, it escapes violation of *DEPEND.¹⁶

Turning to the SV node, *DEPEND in (27) evaluates liquids and nasals by the amount of structure internal to SV that they bear. If we compare (26a) and (26b), liquids have one more feature than nasals— i.e., the feature [approximant]. As a result, in languages where codas are restricted to nasals, coda liquids are always less harmonic than coda nasals. The tableau in (28) illustrates how a nasal is selected over a liquid for an input voiceless obstruent.

(28) input = $C|_{\sigma}$ (coda obstruent)
 $\begin{array}{c} \diagdown \\ L \quad Pl \end{array}$

	*CODALAR	SONNODE	DEPSV	*DEPEND
a. $C _{\sigma}$ $\begin{array}{c} \diagdown \\ L \quad Pl \end{array}$	*			
b. $C _{\sigma}$ $\begin{array}{c} \diagdown \\ SV \quad Pl \end{array}$			*	
c. $C _{\sigma}$ $\begin{array}{c} \diagdown \\ SV \quad Pl \\ \\ [approx] \end{array}$			*	*
d. $C _{\sigma}$ $\begin{array}{c} \diagdown \\ \quad Pl \end{array}$		*		

In the tableau in (28), candidate (a) is a voiceless obstruent, (b) is a nasal, (c) is a liquid, and (d) is an unreleased voiceless obstruent. Candidate (a) loses out first by violating one of the highest ranked constraints, *CODALAR, since it bears a Laryngeal node. (d) also loses out by violating SONNODE since it does not bear any sonority nodes. Between (b) and (c), (c) violates *DEPEND while (b) does not. Therefore, candidate (b) is optimal.

3.2.3 Coda—[ʔ]

As we have seen, one way to resolve a violation of *CODALAR is by changing coda obstruents into sonorants, which bear an SV node, (28). Another way is by under-parsing the Laryngeal node without inserting an SV node. For example, Buginese (a South Sulawesi language) allows only the first part of a geminate, a homorganic nasal and /ʔ/ in coda. When other consonants /r s k/ are in coda, they are reduced to /ʔ/.

(29) Buginese (Rose (1996); originally from Mills (1975))

- a. maʔbinruʔ 'to make'
- b. tikiʔ 'vigilant'
- c. asiŋ 'sarong'
- d. aŋaraŋ 'horse'
- e. mattikirri 'to watch over'
- f. asik-ku 'my sarong'

A similar fact can be observed in Kiowa and Kagoshima Japanese. Kiowa codas can be either sonorants (/m, n, l, y/) or voiceless stops (/p, t/) in careful speech. However, in casual speech, glottal stop replaces /p/ and /t/ in coda position. In Kagoshima Japanese, stops syllabified as codas after word-final high vowel deletion are debuccalized to glottal stop. Examples from both languages are provided below.

(30) Kiowa Coda Obstruents (Watkins, 1984)

	Careful	Casual	
a.	[tʰópkʷæy]	[tʰóʔkyæy]	‘pierce through’
b.	[tʰópkʷ]	[tʰóʔkʷ]	‘shoot / neg’
c.	[tʰátkyé]	[tʰáʔkyé]	‘sever / sg / det’
d.	[bàtpʷ]	[bàʔpʷ]	‘eat / imp (2sg)’

(31) Kagoshima Japanese (Haraguchi 1984:147)

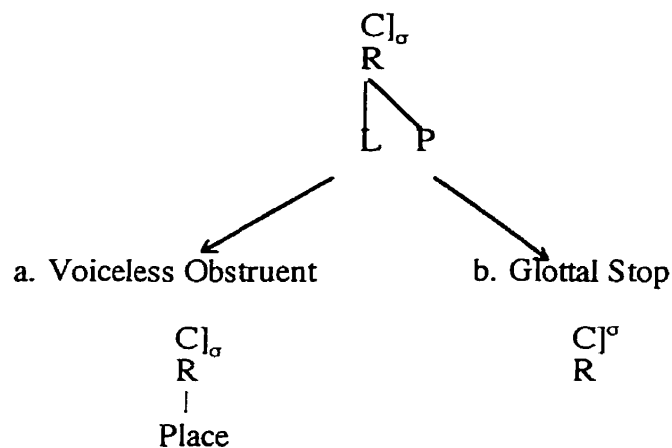
a.	kaki	→	[kaʔ]	‘persimmon’
b.	kagi	→	[kaʔ]	‘key’
c.	tuki	→	[tuʔ]	‘moon’
d.	niku	→	[niʔ]	‘meat’
e.	doku	→	[doʔ]	‘poison’

Languages of this type respect *CODALAR by not having an overt Laryngeal node, but unlike the languages discussed in the preceding section, SV is not inserted to satisfy SONNODE. Thus, as shown earlier in tableau 2 in (23), this type of output is chosen in a language where SONNODE is ranked lower than DEPSV. Under this option, possible coda

consonants are either sonorants which bear SV in the input (therefore, no DEPSV violation) or consonants which bear neither Laryngeal nor SV.

If a segment does not bear either of the sonority nodes as a result of losing Laryngeal, there are two possible outputs: a voiceless (unreleased) obstruent or a glottal stop. The difference between non-Laryngeal voiceless obstruents and glottal stop is that the former have a Place node while the latter does not. See (32) below.

- (32) Laryngeal neutralization:
 Voiceless (unreleased) obstruent vs. Glottal stop
 Voiceless Obstruent (Input)



Among the languages where SONNODE is ranked lower than DEPSV, if *CODAPL ranks lower than MAXPlace which requires Place to be maintained in the output, codas will maintain their place specification and lose only their Laryngeal node. On the other hand, if *CODAPL ranks higher than MAXPlace, codas will lose their place specification as well as their laryngeal specification. Compare the rankings below.

(33) a. MAXPlace » *CODAPL

b. *CODAPL » MAXPlace

Input = C]_σ



(German)			(Buginese)		
	MAXPLACE	*CODAPL		*CODAPL	MAXPLACE
1. C] _σ Pl		*	1. C] _σ Pl	*	
2. C] _σ	*		2. C] _σ		*

In the former type of language, (33a), consonants allowed in coda are either sonorants or Place-bearing voiceless obstruents. Glottal stop is not a possible output in this type of language since the higher ranked constraint, MAXPlace, does not allow this option. In the latter type of language, (33b), sonorants¹⁷ and glottal stops (which result from debuccalization of obstruents) are allowed in coda. Since *CODAPL ranks higher in this type of language, obstruents in coda become glottal stops rather than voiceless obstruents specified for Place.

Although I assume that coda laryngeals that result from debuccalization lack a Laryngeal node, I make no claim about the representations in languages when they exist as underlying segments. See Rose (1996) for recent discussion of various representations.

3.3 Laryngeal Assimilation

Among the languages which have coda Laryngeal neutralization, many of them have Laryngeal assimilation as well. For example, as discussed in Chapter 2, in Ancient Greek, a coda assimilates to the following onset in Laryngeal specification. See the data below.

(34) Ancient Greek (Steriade 1982:231)

a. kleb-dēn	‘stealthy’	:	e-klap-ēn	‘I was cheated’
b. pleg-dēn	‘entwined’	:	plekō	‘to plait’
c. tet ^h lip-tai	‘has been squeezed’	:	t ^h libō	‘to squeeze’
d. lek-teos	‘to be counted’	:	legō	‘to count’
e. strep-tos	‘turned’	:	strep ^h ō	‘to turn’
f. e-dok ^h -t ^h ē	‘it seemed’	:	dok-e-ō	‘to count’

In these data, codas lose their Laryngeal nodes and acquire the Laryngeal specification of the following onset (e-dok^h-t^hē vs. dok-e-ō).

Similarly, in Dutch and Bulgarian, obstruents become voiceless in coda except when they are followed by onset voiced obstruents where they instead become voiced. Data from Dutch are provided in (35).

(35) Dutch (Mascaró 1987, Kenstowicz 1994)

a. hui[z]en	‘houses’	e. a[s]en	‘ashes’
b. hui[s]	‘house’	f. a[s]	‘ash’
c. hui[sk]ammer	‘living room’	g. a[zb]ack	‘ashtray’
d. hui[zb]aas	‘landlord’		

As these examples show, when a voiced obstruent is syllabified as a word-final coda, it becomes devoiced (see (35b) and (35f)). In word-medial position, coda obstruents are voiced when followed by a voiced obstruent (35d,g), and voiceless when followed by a voiceless obstruent (35c).

A similar pattern can be found in Bulgarian, as shown in (36). Voiced obstruents are devoiced in word-final coda position (see (36b)). In word-medial position, they take on the voicing specification of the following onset obstruent.

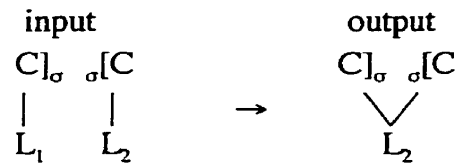
(36) Bulgarian

a. grád-ové		[gradové]	‘cities’	(Scatton 1984)
b. grád		[grát]	‘city’	(Scatton 1984)
c. grád-k-a		[glátka]	‘smooth (fem. sg.)’	(Scatton 1984)
	cf.	[gládŭk]	‘smooth (masc. sg.)’	(Scatton 1984)
d. svát-b-a		[svádba]	‘wedding’	(Scatton 1984)
	cf.	[svátove]	‘matchmakers’	(Scatton 1984)

Notice that in all three languages (Ancient Greek, Dutch and Bulgarian) introduced above, codas acquire the Laryngeal specification of the following onset, whatever it is. Importantly, we cannot analyze these cases as the spreading of sub-Laryngeal features. Such an analysis would require a binary feature system to account for the case where a voiced obstruent become voiceless in front of a voiceless obstruent, or where an aspirated obstruent become unaspirated in front of an unaspirated obstruent. The problem is that the existence of the negative value for laryngeal features cannot be independently supported (Itô & Mester 1986, Lombardi 1991). Moreover, the privativeness of [voice] is supported by other phonological facts addressed in this thesis, in particular, the case of Rendaku discussed in

Chapter 2. Thus, the assimilation phenomena in Ancient Greek, Dutch, and Bulgarian must be analyzed as Laryngeal node assimilation: codas receive Laryngeal from the following onset, as illustrated in (37).

(37) Laryngeal Sharing in Ancient Greek, Dutch, and Bulgarian



The pattern in (37) can be captured through high ranking of both *CODALAR and SONNODE. Recall that *CODALAR prohibits the coda from licensing a Laryngeal node, while SONNODE requires all segments to bear some sonority node. Both of these constraints can be satisfied if the coda receives its sonority node from the following onset.

The difference between languages which exhibit Laryngeal spreading and languages like German, where Laryngeal in coda is neutralized without assimilation, can be attributed to a difference in the ranking between *Crispness* constraints¹⁸ (Itô & Mester 1994) and SONNODE. The syllable Crispness constraint requires syllable edges to be ‘crisp’; in other words, it prohibits the sharing of features or nodes across syllable boundaries, as in the output configuration in (37). See the tableaux below.

(38) Constraints: *CODALAR; SONNODE; CRISP

Input = $C]_{\sigma} \sigma [C$
 $\begin{array}{cc} | & | \\ L & L \end{array}$

Tableau 1: Ancient Greek, Dutch

		*CODALAR	DEPSV	SONNODE	CRISP
	a. $C]_{\sigma} \sigma [C$ $\begin{array}{cc} & \\ L & L \end{array}$	*			
	b. $C]_{\sigma} \sigma [C$ $\begin{array}{c} \\ L \end{array}$			*	
☞	c. $C]_{\sigma} \sigma [C$ $\begin{array}{c} \vee \\ L \end{array}$				*
	d. $C]_{\sigma} \sigma [C$ $\begin{array}{cc} & \\ SV & L \end{array}$		*		

Tableau 2: German:

		*CODALAR	DEPSV	CRISP	SONNODE
	a. $C]_{\sigma} \sigma [C$ $\begin{array}{cc} & \\ L & L \end{array}$	*			
☞	b. $C]_{\sigma} \sigma [C$ $\begin{array}{c} \\ L \end{array}$				*
	c. $C]_{\sigma} \sigma [C$ $\begin{array}{c} \vee \\ L \end{array}$			*	
	d. $C]_{\sigma} \sigma [C$ $\begin{array}{cc} & \\ SV & L \end{array}$		*		

If SONNODE ranks higher than CRISP, as in Ancient Greek and Dutch, the coda will receive Laryngeal from the following onset to be faithful to the former constraint; see candidate (c) in tableau 1 in (38). On the other hand, if the ranking is the reverse as in German, candidate (b) will be selected where the syllable edge is kept crisp by delinking Laryngeal, in spite of the violation of SONNODE.

3.3.1 Final Exceptionality

We have seen that hetero-syllabic Laryngeal assimilation from onset to coda is a consequence of SONNODE. Since SONNODE requires segments to bear either Laryngeal or SV, codas which cannot bear their own Laryngeal node receive Laryngeal from the following onset. This analysis makes an interesting prediction when it comes to word-final position. Although in word-internal clusters, a coda obstruent can receive Laryngeal from the following onset, there is no following onset for a word-final obstruent to assimilate to. There are thus two options: 1. Laryngeal can simply be neutralized in order to satisfy *CODALAR word-finally; or 2. Laryngeal can be maintained, violating *CODALAR in order to satisfy SONNODE. Whether a language selects the former or the latter option is thus determined by the ranking between *CODALAR and SONNODE. Dutch and Ancient Greek are languages which select the former option: *CODALAR » SONNODE. Serbo-Croatian and Hungarian opt for the latter: SONNODE » *CODALAR. Data from Serbo-Croatian and Hungarian are given in (39) and (40) respectively.

(39) Serbo-Croatian (Cho 1990; originally from Grizdic 1969)

a.	rob	‘slave’	e.	drugi	‘different’
b.	ropstavo	‘slavery’	f.	drukâiji	‘more different’
c.	top	‘gun’	g.	svat	‘wedding guest’
d.	tobdžija	‘gunner’	h.	svadba	‘wedding’

(40) Hungarian (Vago 1980)

			-ban/ben ‘in’
a.	kalap	‘hat’	kalab ban
b.	kút	‘well’	kúd ban
c.	zsák	‘sack’	zság ban
d.	lakás	‘apartment’	lakáz ban
			-tól/től ‘from’
e.	rab	‘prisoner’	rap tól
f.	kád	‘tud’	kátt ól
g.	meleg	‘warm’	melekt ől
h.	víz	‘water’	vís t ől

Focussing on Serbo-Croatian, a comparison of (39a) and (39c) shows that the voicing distinction is retained in word-final position. However, the voiced obstruent /b/ in (39a) becomes voiceless when it is followed by voiceless /s/ in (39b). In addition, voiceless /p/ in (39c) becomes voiced when it is followed by the voiced obstruent /dž/ in (39d).

Serbo-Croatian and Hungarian reveal that word-final position sometimes displays different behavior from what is observed word-medially. While *CODALAR forces Laryngeal assimilation in these languages in word-internal clusters, word-finally, obstruents can retain their Laryngeal specification because of SONNODE. Asymmetries of this sort have led some researchers to treat word-final position as special. For example, some languages permit more contrasts in word-final position than in word-internal coda position. This special status of

word-final position has been discussed in the literature under the rubric of *extraprosodicity* (e.g. McCarthy 1979, Itô 1986, Myers 1987). To account for this special status of word-final position, some researchers (Giegrich (1985), Kaye (1990), McCarthy & Prince (1990) and Piggott (1991)) have proposed that extraprosodic consonants are in fact onsets of syllables which contain empty nuclei.¹⁹ However, final consonants in the languages being discussed here do not lend themselves to such an analysis. Crucially, the segments which may appear in word-final position do not differ from the ones which can appear in word-internal coda *except* in their laryngeal features. For example, in Hungarian, any consonant can occupy the coda position except /h/ (John Jensen, personal communication). /h/ is deleted in word-medial coda as well as in word-final position. This fact may suggest that word-final position is not extraprosodic.

In the present analysis of regressive Laryngeal assimilation, final exceptionality is a result of enforcing SONNODE. If both *CODALAR and SONNODE are highly respected together with DEPSV, the result will be Laryngeal assimilation violating CRISP. Since both of these constraints cannot be satisfied word-finally, final exceptionality for coda Laryngeal neutralization should be attested only in those languages where Laryngeal assimilation is observed. Thus, the present analysis predicts that there should be three types of coda Laryngeal neutralizing languages, identified in (41):

(41) Laryngeal Neutralizing Language Types

	Laryngeal Assimilation	Final Exception	Languages
1.	Yes	No	Dutch, Bulgarian
2.	Yes	Yes	Serbo-Croatian, Hungarian
3.	No	No	German
4.	No	Yes	Does Not Exist

Through reranking of three constraints, *CODA_{LAR}, SON_{NODE} and CRISP, there are six possible rankings as can be seen in (42).

(42) Rerankings: *CODA_{LAR}; SON_{NODE}; CRISP

- a. *CODA_{LAR} » SON_{NODE} » CRISP
- b. *CODA_{LAR} » CRISP » SON_{NODE}
- c. CRISP » *CODA_{LAR} » SON_{NODE}
- d. CRISP » SON_{NODE} » *CODA_{LAR}
- e. SON_{NODE} » *CODA_{LAR} » CRISP
- f. SON_{NODE} » CRISP » *CODA_{LAR}

Among the six rankings, two of them, (42d) and (42f), do not command Laryngeal neutralization, so they are not relevant to the present discussion. The tableaux in (43) show that reranking of the three constraints, *CODA_{LAR}, SON_{NODE} and CRISP, yields only the three

types of languages listed in (41), and not the fourth. If the ranking of these three constraints is $*\text{CODA}_{\text{LAR}} \gg \text{SON}_{\text{NODE}} \gg \text{CRISP}$, Tableau 1 in (43), then Laryngeal will be assimilated in word-internal obstruent clusters, and Laryngeal will be neutralized word-finally as in Dutch. If the ranking is $*\text{CODA}_{\text{LAR}} \gg \text{CRISP} \gg \text{SON}_{\text{NODE}}$ or $\text{CRISP} \gg *\text{CODA}_{\text{LAR}} \gg \text{SON}_{\text{NODE}}$, Tableaux 2 and 3, respectively, then Laryngeal will be neutralized both in word-medial obstruent clusters and in word-final position as in German. Finally, if $\text{SON}_{\text{NODE}} \gg *\text{CODA}_{\text{LAR}} \gg \text{CRISP}$, then Laryngeal will assimilate in medial obstruent clusters and Laryngeal will not be neutralized in word-final position as in Serbo-Croatian.

(43) Constraint Reranking: $*\text{CODA}_{\text{LAR}}$; SON_{NODE} ; CRISP

Tableau 1. $*\text{CODA}_{\text{LAR}} \gg \text{SON}_{\text{NODE}} \gg \text{CRISP}$ (Dutch) (= (42a))

Coda Laryngeal
input = $\text{C}|_{\sigma\sigma}[\text{C}$
 | |
 L L

		$*\text{CODA}_{\text{LAR}}$	SON_{NODE}	CRISP
	$\text{C} _{\sigma\sigma}[\text{C}$ L L	*		
	$\text{C} _{\sigma\sigma}[\text{C}$ L		*	
✖	$\text{C} _{\sigma\sigma}[\text{C}$ L			*

Word-final Laryngeal
input = $\text{C}|\#$
 |
 L

		$*\text{CODA}_{\text{LAR}}$	SON_{NODE}	CRISP
	$\text{C} \#$ L	*		
✖	$\text{C} \#$		*	

Tableau 2. *CodaLar » Crisp » SonNode (German) (=42b))

Coda Laryngeal
input = C]_σ σ[C
 | |
 L L

Word-final Laryngeal
input = C]#
 |
 L

		*Coda Lar	Crisp	Son Node
	C] _σ σ[C L L	*		
☞	C] _σ σ[C L			*
	C] _σ σ[C └─┬─┘ L		*	

		*Coda Lar	Son Node	Crisp
	C]# L	*		
☞	C]#		*	

Tableau 3. Crisp » *CodaLar » SonNode (German) (=42c))

Coda Laryngeal
input = C]_σ σ[C
 | |
 L L

Word-final Laryngeal
input = C]#
 |
 L

		Crisp	*Coda Lar	Son Node
	C] _σ σ[C L L		*	
☞	C] _σ σ[C L			*
	C] _σ σ[C └─┬─┘ L	*		

		Crisp	*Coda Lar	Son Node
	C]# L		*	
☞	C]#			*

Tableau 4. SonNode » *CodaLar » Crisp (Serbo-Croatian) (= (42e))

Coda Laryngeal
input = C]_{oo}[C
| |
L L

		SON NODE	*CODA LAR	CRISP
	C] _{oo} [C L L		*	
	C] _{oo} [C L	*		
☞	C] _{oo} [C L			*

Word-final Laryngeal
input = C]#
|
L

		SON NODE	*CODA LAR	CRISP
☞	C]# L		*	
	C]#	*		

As we can see from (42) and (43), the present analysis predicts the absence of one language type: a language which has word-medial coda Laryngeal neutralization (i.e., no assimilation) combined with word-final exceptionality (i.e., voicing contrast maintained). This type of language would require the ranking *CODALAR » SONNODE for medial codas but SONNODE » *CODALAR for final codas. This is a positive effect. According to Cho's (1990a, b) typology of voicing assimilation, there is no language where [voice] is allowed only in word-final codas, and which does not have voicing assimilation. This prediction seems to extend to other Laryngeal features as well. As far as I know, there exists no language which lacks Laryngeal assimilation but has Laryngeal neutralization except word-finally.

At present, it is not clear whether the current analysis of final exceptionality to NOCODA: LARYNGEAL can also account for final exceptionality to NOCODA: PLACE. If a similar analysis held true for Place features, Place final exceptionality would result from *CODAPL and a constraint which required all segments to bear a Place node, PLACENODE. Since Place

final exceptionality would result from the ranking $\text{PLACE} \text{NODE} \gg * \text{CODA} \text{PL}$, there should be no language where debuccalization occurs in word-internal codas (therefore, no assimilation), which suggests the ranking $* \text{CODA} \text{PL} \gg \text{PLACE} \text{NODE}$, but where place features are retained only in word-final position which suggests the ranking $\text{PLACE} \text{NODE} \gg * \text{CODA} \text{PL}$. In other words, if a language restricts place features to word-final codas, then medial codas should assimilate to the following onset in Place to be faithful to both $\text{PLACE} \text{NODE}$ and $* \text{CODA} \text{PL}$; alternatively, a vowel could be epenthesized to avoid syllabifying place-bearing segments as codas. Whether or not such an analysis would hold for Place final exceptionality is beyond the scope of this thesis.

3.4 Summary

In this chapter, we have focussed on two different coda constraints: $\text{NoCODA} : \text{LARYNGEAL}$ ($* \text{CODA} \text{LAR}$) and $\text{NoCODA} : \text{PLACE}$ ($* \text{CODA} \text{PL}$). Under the proposal that Laryngeal (which defines obstruency) and SV (which defines sonorancy) form the class of Sonority nodes, we can provide a straightforward account for several phenomena which are attested across languages: neutralization of laryngeal features, inter-syllabic laryngeal assimilation, coda sonority requirements, and coda debuccalization where obstruents are reduced to glottal stop.

Through the Sonority Node Requirement ($\text{SON} \text{NODE}$) which requires all segments to bear either Laryngeal or SV, we have provided motivation for Laryngeal assimilation in languages where $* \text{CODA} \text{LAR}$ is highly ranked. In addition, $\text{SON} \text{NODE}$ accounts for final exceptionality, where voiced obstruents can only be licensed in word-final coda position.

Through the reranking of *CODALAR and SONNODE, we correctly predict the absence of languages where Laryngeal is neutralized without assimilation in word-medial codas, but where it is not neutralized in word-final position.

—NOTES—

¹ It is irrelevant to my analysis whether or not inputs are fully syllabified. The syllable structure is added in the input for clarity.

² This may follow from Avery & Rice's (1989) Node Activation Convention.

Node Activation Convention

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.

However, it may be the case that Laryngeal is present in a language where there is no laryngeal contrast. I am inclined to think that both Laryngeal and SV are present in all languages since they have an important status as Sonority nodes and, as far as I know, there are no languages which lack sonorant vs. obstruent contrasts.

³ Obstruent clusters are observed more often at word edges. If such clusters can only be found at word-edges, it is likely that the initial (or final) consonant is an adjunct or an appendix.

⁴ One possible way to account for voicing assimilation in onsets would be to refer to phonological government. If the restriction on coda is because the position is governed by the following onset, the second position in the onset should be subject to similar restrictions since it is governed by the preceding consonant. If similar restrictions hold for the second position within an onset, we should find place restrictions which we find between consonants in heterosyllabic coda-onset clusters. As expected, we find voicing agreement in onset clusters. Similarly, we also find cases where the second position in onset is restricted to coronals. For example, in Attic Greek (Steriade 1982), onset clusters are restricted to

“voiceless stop + n, l, or r” and “voiced stop + r, or l” and there are no clusters where the second member is not coronal. However, differently from heterosyllabic clusters, we rarely find place agreement within onset clusters. To determine whether or not the second position in an onset is subject to the same restrictions as a coda, a closer investigation of onset clusters is required. This problem is left open for further study.

⁵ In Beijing Chinese, only /n, ŋ, ɹ/ are allowed in coda (Blevins 1995).

⁶ Japanese also allows the first half of a geminate to be in coda.

⁷ There are also languages where coda fricatives are debuccalized to /h/ while stops are debuccalized to /ʔ/. Kelantan Malay is a language of this type.

Kelantan Malay (Trigo 1991, Rose 1996; originally from Teoh 1988)

	Standard Malay	Kelantan	
a.	ʔasap	ʔasaʔ	‘smoke’
b.	kilat	kilaʔ	‘lightning’
c.	balas	balah	‘finish’
d.	negatef	negatih	‘negative’

In this type of language, the difference between debuccalized /h/ and /ʔ/ seems to be the existence of the feature [continuant].

⁸ Word-final position escapes from the coda restrictions because it is extraprosodic in this language (see McCarthy 1979, Itô 1986, Myers 1987 and Piggott 1991).

⁹ According to my assumption that default SV is nasal, coda stops should become nasals rather than approximants in Hausa. The fact that this is not so may be due to the preservation of place features in the input. If MAXPLACE is highly ranked in Hausa and, as Rice (1996) has suggested, there is a constraint which bans nasals from licensing a Place

node, changing obstruents into approximants would allow Place in the input to be parsed without violating this constraint. In the case of /d/ → [ɾ], coronality is preserved and in /g/ → [w], peripherality is preserved (Rice & Avery (1991) have proposed that Labial and Dorsal are dominated by a node called Peripheral).

¹⁰ The insertion of SV cannot save [voice] since SV is not a proper licenser for this feature. See Section 4.2.2 for discussion.

¹¹ Most of these languages allow the first half of a geminate in coda as well.

¹² Recall from Section 2.1.3 that I adopt Rice's (1993) view that nasal is the default interpretation of SV. Unless there are phonological phenomena which show that [nasal] is active in the phonology of a language, [nasal] will not be projected.

¹³ The original idea of this constraint comes from Heather Goad and Glyne Piggott.

¹⁴ This term refers to intermediate nodes in the tree which dominate terminal features: SV, L, and Place (cf. Rice & Avery 1991). Given that Root dominates all nodes and their dependent features, it has a different status from other organizing nodes.

¹⁵ The *STRUC constraint for syllable structure is proposed by Zoll (1996).

¹⁶ Since Optimality Theory allows constraints to be ranked differently across languages, we should expect to find a language where *DEPEND is undominated. Focussing on place, in such a language, we would only find coronal segments, and not labials or velars. However, such a language is not attested. Moreover, languages always seem to have at least two place contrasts. The absence of a language which does not have place contrasts is no doubt due to some deeper principle of human language which requires languages to have enough contrasts among consonants to maintain a large vocabulary.

¹⁷ The velar nasal also results from debuccalization (Trigo 1988) when the Place node is lost (by respecting *CODAPLACE) with a bare SV remaining in coda. In such a case, both *CODALAR and SONNODE are respected.

¹⁸ The same effect can be obtained by FILLLINE (Itô, Mester & Padgett 1995), a constraint which prohibits the insertion of an association line.

¹⁹ Kaye's (1990) Coda Licensing requires a coda to be licensed by a following onset. For him, therefore, all word-final codas are onsets of empty-headed syllables. Piggott (1991), on the other hand, has proposed this as a parameter. Therefore, he permits languages which have word-final codas.

4 . VOICING ASSIMILATION AND SONORITY

In Chapter 3, we discussed languages like Dutch where hetero-syllabic obstruent clusters agree in voicing. Such examples were analyzed as Laryngeal sharing to satisfy both SONNODE and *CODALAR. In this chapter, we will investigate two other types of voicing assimilation. First, in Section 4.1, we examine intersyllabic regressive voicing assimilation in languages such as Ukrainian. Although the outcome of this type of assimilation looks similar to that of Dutch, I argue that a different account is required for this process. As discussed in the previous chapter, Dutch assimilation can be explained as Laryngeal sharing since coda-onset obstruent sequences completely agree in voicing. In contrast, in Ukrainian, coda obstruents become voiced when the following onset is a voiced obstruent. However, voiced obstruents in coda do not devoice when they are followed by voiceless obstruents. I propose that the regressive voicing assimilation represented by the Ukrainian examples can be accounted for by the SONORITY constraint which was proposed in Section 2.4.

In Section 4.2, we will consider some cases of post-sonorant voicing. The most commonly observed phenomenon of this type is post-nasal voicing, a process which has recently been of issue in the phonological literature (see Itô & Mester 1986; Rice 1993; Itô, Mester & Padgett 1995; Kawasaki 1996; Pater 1996, among others). Post-sonorant voicing

is problematic because the feature [voice] for sonorants has generally been assumed to be unspecified. In this section, I will review some proposals put forward in the literature to deal with this problem. I will then examine some asymmetries between post-nasal voicing and other types of post-sonorant voicing.

Finally, in Section 4.2.5, I will compare post-sonorant voicing with pre-sonorant voicing, and argue that these two types of assimilation processes must be treated differently.

4.1 Regressive Voicing Assimilation

Recall from Section 3.3 that hetero-syllabic obstruent clusters agree in voicing in Dutch. In (1c, d), it can be seen that coda obstruents acquire the Laryngeal specification of the following onset.

(1) Dutch		(Mascaró 1987, Kienstowicz 1994)	
a. hui[z]en	‘houses’	e. a[s]en	‘ashes’
b. hui[s]	‘house’	f. a[s]	‘ash’
c. hui[sk]ammer	‘living room’	g. a[zb]ack	‘ashtray’
d. hui[zb]aas	‘landlord’		

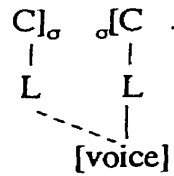
Although it has some similarity to the Dutch pattern, regressive voicing assimilation in Ukrainian exhibits a different pattern, as reflected in (2) below.

(2) Ukrainian (data are from Humesky (1980) and Cho (1990a))¹

a. /borot'bá/	→	[borod'bá]	'battle'
/vokzái/	→	[vogzái]	'train station'
/ják že/	→	[jáɡ že]	'maybe; as if'
b. bereg			'shore'
viz			'cart'
golub			'pigeon'
c. rid-ko			'rare'
xobty			'truck gen. sg.'
viez-ty			'to drive'

The examples in (2a) show that coda voiceless obstruents become voiced when the following onset is a voiced obstruent. From these data alone, this phenomenon looks very similar to that of Dutch. However, as we can see in (2b), Ukrainian does not have coda devoicing. In addition, when the following onset is voiceless, the coda does not take on its 'voicelessness' as in (2c). These data reveal that the Laryngeal sharing analysis for Dutch and Ancient Greek discussed in Section 3.3 cannot be applied to Ukrainian. If Laryngeal is being shared, voicelessness as well as voicing should be acquired from onsets. Instead, in Ukrainian voicing assimilation, only the feature [voice] spreads, not the entire Laryngeal node. This is illustrated in (3).

(3) Ukrainian Voicing Assimilation



From a comparison of Ukrainian and Dutch, we can conclude that the motivation for the two assimilation processes cannot be the same. In Chapter 3, I proposed that assimilation in Dutch is due to high ranking of *CODALAR and SONNODE. Since *CODALAR does not allow Laryngeal to be present in coda, the coda receives its Laryngeal node from the following onset. I suggest that the Ukrainian case is due to the SONORITY constraint introduced in Section 2.4. Recall that in a hetero-syllabic sequence C_1C_2 , C_2 must be less than or equal in sonority to C_1 . This constraint, which prohibits a hetero-syllabic sequence where the onset is more sonorous than the preceding coda, is repeated in (4).

(4) Syllable Contact Constraint (SONORITY)²

In a hetero-syllabic consonant cluster, $C_1]_{\sigma}]_{\sigma} C_2$, the sonority value of C_2 ($\{C_2\}$) minus the sonority value of C_1 ($\{C_1\}$) cannot be more than 0.

Violations of this constraint are calculated in a gradient manner. For example, if the sonority value of the coda is 3 and that of the following onset is 1, the sequence violates SONORITY by 2 (this is indicated in the tableaux by two asterisks).

In Ukrainian, assimilation only takes place when a voiceless obstruent in coda is followed by a voiced obstruent in onset. Recall from Chapter 2 that [voice] is dependent on

both sonority nodes and that it figures into the calculation of sonority. Since [voice] adds a value of one, and the Laryngeal node adds no value, a voiced obstruent has a sonority value of 1 while a voiceless obstruent has a value of 0. Therefore, a voiceless + voiced sequence in a heterosyllabic obstruent cluster violates SONORITY by 1 (and acquires one asterisk).

If SONORITY in (4) is ranked higher than CRISP which, in requiring syllable edges to be ‘crisp’, prohibits spreading, voiceless + voiced obstruent sequences will be repaired as voiced + voiced in order to be faithful to SONORITY. See the tableau in (5).

(5) Ukrainian Voicing Assimilation: SONORITY » CRISP

input =
$$\begin{array}{cc} C]_{\sigma} & _{\sigma} [C \\ | & | \\ L & L \\ & | \\ & [vce] \end{array}$$

		SONORITY	CRISP
	a. $\begin{array}{cc} C]_{\sigma} & _{\sigma} [C \\ & \\ L & L \\ & \\ & [vce] \end{array}$	*	
☞	b. $\begin{array}{cc} C]_{\sigma} & _{\sigma} [C \\ & \\ L & L \\ \diagdown & / \\ & [vce] \end{array}$		*

There are other possible ways to resolve the violation of SONORITY present in the input in (5). One option would be to fuse Laryngeal. As a result, coda and onset would share one Laryngeal node. This option results in a violation of LINEARITY (McCarthy & Prince 1996) which requires any precedence relationship in the input to be maintained in the output. Since

the precedence relationship between the Laryngeal node in coda and the Laryngeal node in the following onset is not maintained in the output where these two Laryngeals are fused, LINEARITY is violated in this representation. Although (5b) violates *CODALAR, it satisfies LINEARITY. In reality, the ranking between *CODALAR and LINEARITY cannot be determined from the data given in (2). Thus, we cannot tell which representation should be chosen for Ukrainian outputs. However, the important point is that [voice] sharing or Laryngeal fusion is caused by the high ranking of SONORITY.

Another option which would resolve the SONORITY violation in Ukrainian would be to delink the feature [voice] from the onset. However, languages tend to parse onset features more faithfully than coda features. This may be related to the fact that languages tend not to restrict what can appear in an onset. If we were to encode such a tendency in the grammar in terms of a constraint, it would be along the lines of the constraint in (6).

(6) MAXONSET (Consonantal) (MAXONS)³

Consonantal features must be parsed in onset.

This constraint expresses the fact that languages syllabify consonants maximally into onset as long as the result does not violate syllable wellformedness restrictions. In other words, this constraint means that onsets prefer to be consonants. Similar constraints which require consonantal point of articulation to be licensed by an onset have been proposed by Steriade (1995). For an extension of Steriade's proposal, see Humbert (1996). Consonantal features would include consonantal place (C-Place) (Clements & Hume 1995), continuancy, and Laryngeal features in contrast to SV features which are more vowel-like. In Ukrainian,

MAXONS is ranked higher than CRISP; therefore, the resolution of SONORITY by delinking [voice] will not be selected as optimal.

4.2 Post-Sonorant Voicing

4.2.1 Voicing Paradox

As discussed in Section 2.1.1.4, much research conducted within the framework of Underspecification Theory has shown that [voice] for sonorants is phonologically inactive. Many languages, however, have voicing assimilation triggered by sonorants. One of the most commonly observed voicing assimilation processes is post-nasal voicing where obstruents immediately preceded by a nasal become voiced. Some languages which exhibit this process are Luyua (Herbert 1986), Kikuyu, Vai (Welmers 1976), OshiKwanyama (Steinbergs 1985), Ijò (Williamson 1987), Yamato-Japanese, and Zoque (Wonderly 1951). The data in (7) come from Zoque.

(7) Zoque Post-Nasal Voicing (Wonderly 1951)

- a. N- '1st sg. poss.+ stem
- | | | | | | |
|---|---|--------|---|---------|----------|
| N | + | tatah | → | ndatah | 'father' |
| N | + | kwarto | → | ŋgwarto | 'room' |
| N | + | plato | → | mblato | 'plate' |
- b.
- | | |
|-----------|---------------------|
| camba | 'he speaks' |
| nʌmgeʔtu | 'he also said' |
| hañʌkʌmis | 'you did not do it' |
| maŋgeʔtu | 'he also went' |
| šuhpuñbuñ | 'soapberry' |

Zoque does not allow nasal + voiceless stop clusters.⁴ As a result, nasals trigger voicing of the following obstruent. Thus, contrary to the widely-accepted view that [voice] is redundant and therefore unspecified for sonorants, a voicing feature seems to be required for nasals in order to account for post-nasal voicing.

4.2.2 Licensing Cancellation

As a solution to the problem that [voice] is redundant for sonorants and yet active in processes like post-nasal voicing, Itô, Mester & Padgett (1995) offer a constraint-based account within the framework of Optimality Theory.⁵ In the spirit of underspecification theory, they formulate the universally undominated constraint in (8).

(8) Licensing Cancellation: If $F \supset G$, then $\neg(F \wedge G)$.

“If the specification F implies the specification G , then it is not the case that F licenses G .” (Itô, Mester & Padgett 1995:580)

It follows from this constraint that sonorants cannot license [voice] since this feature is redundant for sonorants. However, another constraint $\text{SONVOI} \left(([\text{sonorant}] \supset [\text{voice}]) \right)$ expresses the fact that sonorants are fundamentally voiced segments and it therefore requires sonorants to have [voice]. Since Licensing Cancellation does not allow sonorants to license [voice] by themselves, a representation where the feature [voice] is linked to a sonorant is well-formed if and only if a proper licenser for [voice], namely an obstruent, is also present. See the representations below.

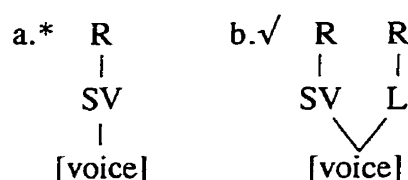
- (9) a. [voice] is not licensed b. [voice] is licensed
- $$\begin{array}{c} * N]_{\sigma} \text{ }_{\sigma} [C \\ | \\ [\text{voice}] \end{array}$$

$$\begin{array}{c} N]_{\sigma} \text{ }_{\sigma} [C \\ \diagdown \quad \diagup \\ [\text{voice}] \end{array}$$

In (9a), [voice] is not licensed since $N(\text{asal})$ cannot be a licenser for [voice] according to Licensing Cancellation. In (9b), on the other hand, [voice] is licensed, not by the nasal, but by the onset obstruent which is a proper licenser for [voice]. It should be noted that, crosslinguistically, post-nasal voicing (and post-sonorant voicing in general) occurs only under specific conditions: between a coda and the following onset where the coda is nasal as in (9b), or in complex segments.⁶ The analysis of post-nasal voicing presented so far is an extension of the proposal by Itô, Mester & Padgett (1995).

I adopt Itô, Mester & Padgett's (1995) universally undominated constraint Licensing Cancellation from (8). I will now delineate the consequences which result from the incorporation of this constraint into a theory that includes feature geometry. Recall from Chapter 2 that I adopt a geometry where the feature [voice] is dominated by both sonority nodes, SV and Laryngeal. Licensing Cancellation prohibits the feature [voice] from being licensed by the SV node because being sonorant, which means being specified for SV, implies being spontaneously voiced. Since I assume that all features must be licensed⁷ to be phonetically realized (parsed or 'linked'), SV by itself cannot sanction the parsing of [voice]. However, [voice] can be parsed by SV when [voice] is parasitically licensed by a Laryngeal node which is present in the following obstruent as in (10b).

(10) Parasitic Licensing of [voice]

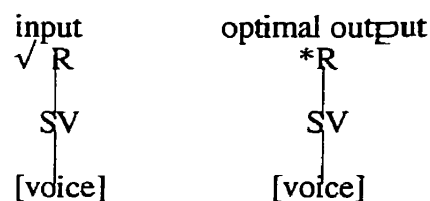


Since Licensing Cancellation is universally undominated, no language will allow (10a) as an optimal surface representation.

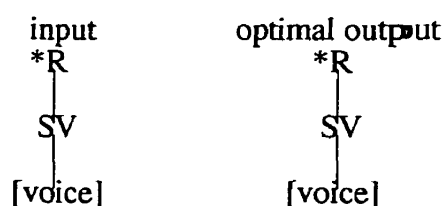
A constraint which is universally undominated is crucially different from one which has the status of a principle of Universal Grammar. A principle of Universal Grammar governs the wellformedness of both input and output. On the other hand, a constraint which is universally undominated does not constrain the input. To illustrate, if we assume that Licensing Cancellation is undominated, this allows the representation in (11a) to be a possible input although it will never be selected as an optimal output in any language. If we

assume that Licensing Cancellation is a UG principle, on the other hand, this representation will never be a possible input; see (11b).⁸

(11) a. Licensing Cancellation as universally undominated



b. Licensing Cancellation as a principle of UG



Since I assume that all features must be licensed to be parsed, the output representations in (11) are not possible in any language, regardless of the input. However, the option in (11a) has empirical advantages over that in (11b); these will be discussed in detail in Chapter 5.

Although some researchers have assumed that there are universally undominated constraints, how these constraints are to be encoded in Optimality Theory is not clear. Universally undominated constraints and constraints which are undominated in a particular language will both never be violated in that language; therefore, the ranking among these two types of constraints in a given language can not be empirically determined.

4.2.3 Licensing Cancellation and Parasitic Licensing

So far, we have discussed a case of parasitic licensing driven by Licensing Cancellation. The representations in (12) reflect my proposal regarding the presence or absence of [voice] thus far.

(12) Output Representations

a. sonorant ⁹	b. obstruent (voiceless)	c. obstruent (voiced)	d. sonorant + obstruent ("Parasitic Licensing")
R	R	R	R R
SV	L	L	SV L
			└─┘
		[vce]	[vce]
e. *sonorant			
R			
SV			
[vce]			

As discussed earlier, (12e) is not a possible optimal output given that Licensing Cancellation is universally undominated. A comparison of the representations in (12d) and (12e) shows that [voice] for sonorants can only be licensed when an obstruent, which is a proper licenser of [voice], follows.

The tableaux in (13), (14) and (15) show how Licensing Cancellation (LC) accounts for both the voicing underspecification effects and post-sonorant voicing. First, consider the tableau in (13) which illustrates how a candidate in which [voice] for a sonorant is unspecified is selected when it is not followed by an obstruent.

(13) Nasals: Not followed by an obstruent

input = $[N V]_{\sigma}$
 |
 SV
 |
 [vce]

		LC	MAXVce
	a. $[N V]_{\sigma}$ SV [vce]	*	
13	b. $[N V]_{\sigma}$ SV		*

In the tableau in (13), candidate (a), which is identical to the input, loses over (b) because it crucially violates Licensing Cancellation. Therefore, although (b) violates MAXVCE, which requires the feature [voice] to be parsed, it is selected as optimal. Since Licensing Cancellation is universally undominated, the ranking between it and MAXVCE is inalterable. In other words, there will be no language in which MAXVCE is ranked higher than Licensing Cancellation. Therefore, no language will choose (a) as the optimal output with the result that onset nasals should never behave as voiced.¹⁰ To my knowledge, this is empirically supported.

Consider next the case where sonorants are followed by an obstruent which can serve as a licenser for [voice].

(14) Nasal + Obstruent

LC » MAXVCE » CRISP

input = N]_σ α[C
 | |
 SV L
 |
 [vce]

		LC	MAX VCE	CRISP
	a. N] _σ α[C SV L [vce]	*		
US	b. N] _σ α[C SV L / [vce]			*
	c. N] _σ α[C SV L		*	

(15) Nasal + Obstruent

LC » CRISP » MAXVCE

input = N]_σ α[C
 | |
 SV L
 |
 [vce]

		LC	CRISP	MAX VCE
	a. N] _σ α[C SV L [vce]	*		
	b. N] _σ α[C SV L / [vce]		*	
US	c. N] _σ α[C SV L			*

Among the three constraints in (14) and (15), only the ranking between MAXVCE and CRISP is variable. Thus the (a) candidates will not be selected in any language. In the (b) candidates, [voice], which is linked to SV in the input, is parasitically licensed by the Laryngeal node of the following onset. In other words, MAXVCE is respected at the cost of violating CRISP.¹¹ This situation is reversed in the (c) candidates where input [voice] is unparsed. In this candidate, CRISP is satisfied by not parsing [voice], although this induces a violation of MAXVCE. In languages which select the ranking in (14), where MAXVCE is above CRISP, post-nasal voicing will result. In languages which select the ranking in (15), where coda [voice] is underparsed, coda sonorants will not behave like voiced segments.

We can conclude from the three tableaux in (13)-(15) that Licensing Cancellation demands [voice] for sonorants to be unparsed unless it is linked to Laryngeal. On the other hand, coda sonorants have two choices: 1) they can behave as voiced through parasitic-licensing, or 2) they can behave as voiceless through underparsing of [voice] due to Licensing Cancellation. Thus, the combination of Licensing Cancellation with the model of subsegmental structure proposed here predicts that sonorants will not behave as voiced in a representation where [voice] is not linked up to Laryngeal. This, in fact, is where we observe voicing underspecification effects. On the other hand, since [voice] can be present for sonorants when it is linked to Laryngeal, this is the only context where we find nasals behaving as voiced. This issue will be taken up again in Chapter 5.

Thus, Itô, Mester & Padgett's analysis and the one presented here account for post-nasal voicing in the same way. One might wonder then why the introduction of SV and Laryngeal structure is necessary. In the following sections, I will demonstrate that the structural account proposed here provides the necessary tools to account for the full range of facts about voicing assimilation triggered by sonorants. Itô, Mester & Padgett's analysis fails to capture these facts in a straightforward manner.

4.2.4 Asymmetry in Post-Sonorant Voicing

4.2.4.1 NoLink

Although Licensing Cancellation accounts for post-sonorant voicing and for the underspecification of [voice] in sonorants, one problem still remains. In post-sonorant voicing, an asymmetry is observed across languages. The cross-linguistic tendency seems to

be for languages to have post-nasal voicing without post-approximant voicing; e.g., Zoque (Wonderly 1951), Kikuyu (Davy & Nurse 1982, Lombardi 1991). If a language has post-approximant voicing, then it seems that it will also have post-nasal voicing; e.g., Basque (Hualde 1988). Examples from Zoque and Basque¹² which demonstrate this difference are provided below.

(16) Zoque (Wonderly 1951)

a. Nasal + C

camba	'he speaks'
naŋgeʔtu	'he also said'
hañʔakamis	'you did not do it'
maŋgeʔtu	'he also went'
ʃuhpuñbuñ	'soapberry'

b. Approximant + C

flawta	'harmonica'	*flawda
kuytʔam	'avocado'	*kuydʔam
kwerpo	'baby'	*kwerbo
porke	'because'	*porge

(17) Basque (Hualde 1988: 231-232)

a. Nasal + C

lan-tu	[lan <u>du</u>]	'labot'	(perfective)
ken-tu	[ken <u>du</u>]	'take away'	(perfective)
egin-ko	[egin <u>go</u>]	'do, make'	(future)
eśan-ta	[eśan <u>da</u>]	'said'	(participial)

b. Lateral + C

afal-tu	[afa l du]	'have dinner'	(perfective)
il-ko	[il g o]	'die, kill'	(future)
il-ta	[il d a]	'dead'	(participial)

To account for this asymmetry, Itô, Mester & Padgett (1995) propose the NoLINK family of constraints which militates against sonorant-obstruent linkages. See (18) below (where V=vowel, G=glide, L=liquid, N=nasal, C=Obstruent).

(18) Constraint family: NoLINK

NO-VC-LINK » NO-GC-LINK » NO-LC-LINK » NO-NC-LINK

Constraints within the NoLINK family are ranked 'intrinsically and universally'; in other words, the order within the ranking is inalterable. With this family of constraints, Itô, Mester & Padgett (1995) can capture the following crosslinguistic tendency: nasal-obstruent linkage (with [voice], for example) is more common than such a linkage between a liquid/glide/vowel and following obstruent. By combining this constraint family with SONVOI, which requires sonorants to have [voice], Itô, Mester & Padgett can account for the fact that post-nasal voicing is less marked than post-approximant voicing. The ranking in (19) holds for languages in which, among sonorants, only nasals trigger voicing assimilation. SONVOI must intervene between NO-NC-LINK and the remaining NoLINK constraints.

(19) Post-Nasal Voicing

		NO-LC-LINK	SONVOI	NO-NC-LINK
a.	1. N C \ / [vce]			*
	2. N C		*	
b.	1. L C \ / [vce]	*		
	2. L C		*	

Given Licensing Cancellation, inserting [voice]¹³ to satisfy SONVOI always forces a sonorant-obstruent linkage, resulting in a violation of one of the NO LINK constraints. Since the first three NO LINK constraints (NO-VC-LINK, NO-GC-LINK, NO-LC-LINK) rank higher than SONVOI, SONVOI must be violated in order to satisfy these constraints. As a result, there will be no [voice] insertion for vowels, glides or liquids; see (19b). However, since SONVOI ranks higher than NO-NC-LINK, the former must be satisfied at the expense of violating the latter. The result is post-nasal voicing; see (19a).

Itô, Mester & Padgett can account for the asymmetry of post-sonorant voicing assimilation through the inalterable ranking among members of the NO LINK family. However, the NO LINK family of constraints is no more than a restatement of the fact that post-nasal voicing is more common than other types of post-sonorant voicing. Secondly, related to this, the constraints are not independently motivated. Thirdly, the NO LINK family does not take syllable structure or the direction of assimilation into consideration, in spite of the fact that parasitic-licensing of [voice] is most often attested from onset to preceding coda

and not in the reverse configuration. In order to handle this restriction, a theory of licensing which is different from the one Itô, Mester & Padgett propose is required; this issue will be discussed in Section 4.3.1. Finally, their analysis encounters an empirical problem when we consider some facts from Yamato-Japanese. In Yamato-Japanese, obstruents become voiced after glides in some cases and do not in others. This problem will be discussed in Chapter 5.

4.2.4.2 *NC̥

Pater (1996) has taken a different approach to the asymmetry observed in post-sonorant voicing. He proposes a constraint which indicates that nasal + voiceless obstruent sequences are marked, *NC̥.

(20) *NC̥

No nasal + voiceless obstruent sequences

This phonetically-grounded constraint is supposed to express the difficulty that speakers have in producing a sequence of nasal + voiceless obstruent. Pater (1996) cites Huffman's (1993:310) observation that "the raising of the velum occurs very gradually during a voiced stop following a nasal segment, with nasal airflow only returning to a value typical of plain obstruents during the release phase" and argues that a nasal + voiced obstruent sequence allows "a more leisurely raising of the velum than an NC" (NC=nasal + voiceless obstruent). Post-nasal voicing is one strategy which languages use to avoid sequences which

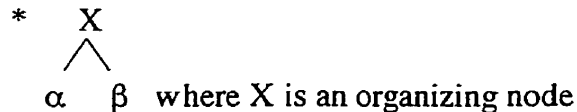
would result in a violation of $*N\underset{\circ}{C}$. Under such an analysis, post-nasal voicing is not necessarily an assimilation process. Voicing on the obstruent does not have to come from the preceding nasal; rather, it may be ‘inserted’ to avoid yielding a nasal + voiceless obstruent sequence.

Since the $*N\underset{\circ}{C}$ analysis does not require nasals to be specified for [voice], the paradox observed between voicing underspecification and voicing specification in NC clusters does not arise. Nevertheless, the analysis cannot be extended to other post-sonorant voicing phenomena. As shown in Section 4.2.4, sonorant consonants other than nasals also induce voicing of the following obstruent. For example, Basque has both post-nasal and post-approximant voicing as observed earlier in (17). Pater’s $*N\underset{\circ}{C}$ analysis requires us to treat post-nasal voicing and post-approximant voicing differently. Consequently, this approach does not straightforwardly extend to the asymmetry observed between post-nasal voicing and post-approximant voicing. As mentioned earlier, post-approximant voicing is always accompanied by post-nasal voicing and not the other way around. Even if we propose another constraint, $*L\underset{\circ}{C}$, which bans approximant + voiceless obstruent sequences to account for post-approximant voicing, the problem of the asymmetry will remain unresolved. The only way to account for it would be to assume that $*N\underset{\circ}{C}$ is always ranked higher than $*L\underset{\circ}{C}$. An additional complication is that $*N\underset{\circ}{C}$ is proposed as a phonetically-grounded constraint, but, it is not clear if there is any phonetic motivation for $*L\underset{\circ}{C}$. This being the case, we would be no further ahead than with the NOLINK family of constraints.

Under my analysis, the effect of Itô, Mester & Padgett's (1995) SONVOI ([sonorant] $\supset$ [voice]) is captured through the combination of MAX [voice] and the assumption that sonorants have [voice] in the input. By themselves, neither SONVOI nor MAX [voice] can capture the fact that, among sonorants, only nasals spread voicing in many languages. The problem is that these constraints do not distinguish between types of sonorants for the presence or absence of a [voice] specification. Recall that Itô, Mester & Padgett (1995) rely on the NOLINK family of constraints to solve this problem. However, I have argued in Section 4.2.4.1 that their account is problematic.

My alternative analysis of the asymmetry adopts the segmental version of *COMPLEX from Padgett (1995) and Goad (1996, in press). This constraint is a member of the No Complex Structure constraint family together with *DEPEND which was introduced in Section 3.2.2. See (21).

(21) *COMPLEX



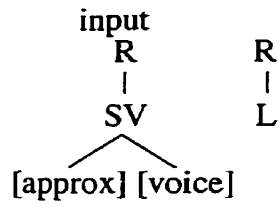
*COMPLEX was originally proposed by Prince & Smolensky (1993) as a constraint which prevents more than one segment from being associated to a single syllable position. The version of *COMPLEX here is an extension of their constraint to segmental structure.

*COMPLEX is required to capture the fact that complex segments such as [labial, dorsal]_{place} are distributionally more marked than their non-complex counterparts, either [labial]_{place} or [dorsal]_{place}. The same situation seems to hold for Laryngeal dependents. For example, voiced aspirated consonants, [SG, voice]_L, are more marked than voiceless ones, [SG]_L.

Recall that I assume that [voice] is present in the input for all sonorants. MAX [voice] prefers [voice] to be parsed in the output. However, sonorant consonants other than nasals have the feature [approximant] underlyingly (see (22)). Maintaining the underlying [voice] specification on these segments in the output would therefore result in a violation of *COMPLEX, (22b). I propose that *COMPLEX is higher-ranked than MAX [voice] in languages where only nasals spread voicing to the following obstruent.

(22) Approximant (+ obstruent)

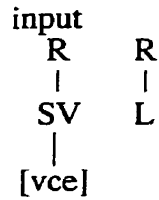
Ranking: MAXAPPROX., *COMPLEX » MAXVCE



		MAXAPPROX	*COMPLEX	MAXVCE
a.	<pre> R R SV L [app] </pre>			*
b.	<pre> R R SV L / \ [app][vce] </pre>		*	

Since nasals are unmarked sonorants, recall that in most languages the feature [nasal] will not be specified in the input. [voice] can thus be maintained in the output representation of a nasal without incurring a violation of *COMPLEX.¹⁴ Consequently, the representation with post-nasal voicing in (23b) is selected over (23a).

(23) Nasals (+ obstruent)



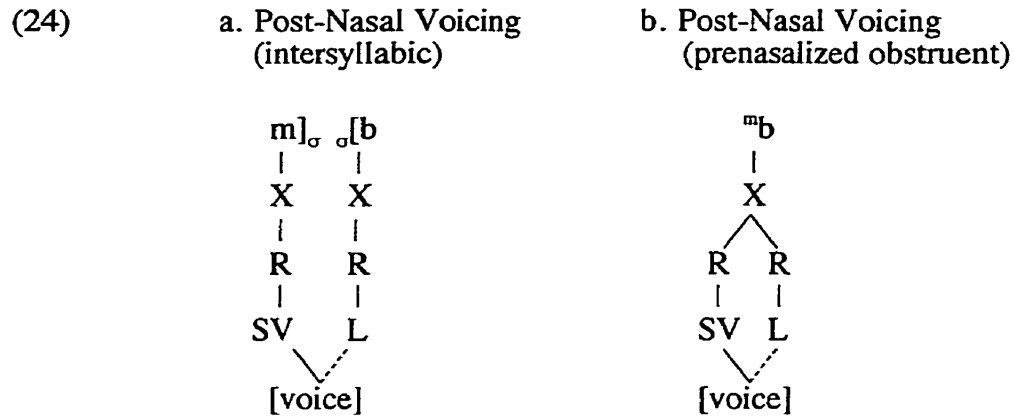
Ranking: *COMPLEX » MAX [voice]

		*COMPLEX	MAXVCE
	a. <pre> R R SV L </pre>		*
☞	b. <pre> R R SV L \ / [vce] </pre>		

The constraint *COMPLEX, in combination with the underlying representation of nasals, thus accounts for why, among sonorants, only nasals spread voicing in many languages.¹⁵ In summary, we can account for three types of languages with *COMPLEX. Among the three constraints, *COMPLEX, MAXVCE, and CRISP, if CRISP outranks MAXVCE, no post-sonorant voicing of any type results as in English. If the ranking is MAXVCE » *COMPLEX, CRISP, we would expect post-sonorant voicing as in Basque. If the ranking is *COMPLEX » MAXVCE » CRISP, the result is post-nasal voicing as in Zoque.

4.2.5 Prenasalized Segments

So far, we have discussed post-nasal voicing in intersyllabic contexts. However, there is another configuration where [voice] for nasals can be parasitically licensed. The two possibilities are illustrated in (24).



In the representations above, place agreement has been omitted. The dotted line indicates an association which may or may not be specified in the input. The representation in (24a) is a nasal + voiced obstruent sequence and (24b) is one representations for a prenasalized voiced obstruent.¹⁶ In both representations, [voice] in SV is parasitically licensed by the following Laryngeal node, which is a proper licenser of [voice].

As shown in (24b), the nasal and oral parts of a pre-nasalized segment are dominated by a *single* position. This point will become crucial when the possible configurations for parasitic licensing are defined in Section 4.3.1.

Returning to the representation in (24b), [voice] which is a property of the nasal part of a prenasalized obstruent can be parasitically licensed by Laryngeal which is a property of the oral part. Thus, according to the parasitic licensing analysis presented so far, voicing agreement in prenasalized obstruents should be common, just as is intersyllabic post-nasal voicing. Consistent with this prediction, post-nasal voicing in prenasalized segments is widely attested across languages. Some examples from Kikuyu (Armstrong 1967; Mugane & Gerfen 1993) and Ndali (Vali 1972) are given below.

(25) Kikuyu (Mugane & Gerfen 1993)

a. tɛma	'cut'	N + tɛma	→	ⁿ dɛma	'cut me'
b. šona	'lick'	N + šona	→	ⁿ ʃona	'lick me'
c. kona	'hit'	N + kona	→	^ŋ gona	'hit me'

(26) Ndali (Vail 1972)

a. /iN + puno/	→ [i ^m buno]	'nose'
b. /iN + tunye/	→ [i ⁿ dunye]	'banana'
c. /iN + kunda/	→ [i ^ŋ gunda]	'dove'

As we can see, in Kikuyu and Ndali, as well as in other Bantu languages, prenasalization of obstruents yields segments which are voiced. In other languages, such as Zande, South Gomen and Fijian (Herbert 1986), all underlying prenasalised stops are voiced.

To summarize, voicing agreement is a very common phenomenon, not only in heterosyllabic NC clusters, but also in prenasalized obstruents. In both cases, voicing is

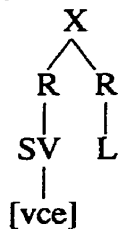
accounted for through a combination of Licensing Cancellation and the subsegmental structure proposed in this thesis.

Before we conclude this section, we will address one final issue. According to Herbert's (1986) observations, prenasalized voiced obstruents are crosslinguistically more common than their voiceless counterparts: if a language has prenasalized voiceless obstruents, it also has voiced ones. To my knowledge, none of the analyses which have been proposed in the literature can formally account for this entailment relation. Since the present analysis adopts the optimality-theoretic premise of free reranking of constraints, it also fails to capture the favored status of prenasalized voiced obstruents.

The preference for voicing assimilation in prenasalized segments can be accounted for by Licensing Cancellation in combination with MAXVCE. However, there must also be a constraint which discourages voicing assimilation since there are some languages which have prenasalized voiceless obstruents as well (e.g., Ganda, Rundi (Herbert 1986)). One potential candidate is DEPLINK which militates against the insertion of association lines (Pulleyblank 1994).

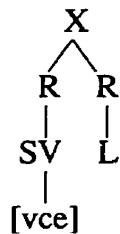
(27) Post-Nasal Voicing in Prenasalized Obstruents

1. input = ⁿt



		MAXVCE	DEPLINK
	1. ⁿ t <pre> graph TD X --> R1[R] X --> R2[R] R1 --> SV[SV] R1 --> L[L] SV --> vce1["[vce]"] </pre>	*	
☞	2. ⁿ d <pre> graph TD X --> R1[R] X --> R2[R] R1 --> SV[SV] R1 --> L[L] SV --> vce2["[vce]"] </pre>		*

2. input = ⁿt



		DEPLINK	MAXVCE
☞	1. ⁿ t <pre> graph TD X --> R1[R] X --> R2[R] R1 --> SV[SV] R1 --> L[L] SV --> vce4["[vce]"] </pre>		*
	2. ⁿ d <pre> graph TD X --> R1[R] X --> R2[R] R1 --> SV[SV] R1 --> L[L] SV --> vce5["[vce]"] </pre>	*	

DEPLINK would be dominated by MAXVCE in a language where all prenasalized segments are voiced. If the ranking is reversed, however, all prenasalized obstruents should be voiceless and such a language does not exist. Thus, this analysis is no further ahead in accounting for the entailment relation. Moreover, the distributional fact about prenasalized segments makes us suspicious about an analysis which relies on a constraint such as DEPLINK which militates against the addition of an association line.

A more fruitful direction in which to look for an explanation of the distributional markedness of prenasalised segments would be to propose a constraint which requires contrasts to be maintained (KEEPCONTRAST) and to only opt for constraints which motivate voicing assimilation and none which militate against assimilation in prenasalized segments. If a language has both prenasalized voiced and voiceless segments and the ranking is MAXVCE » KEEPCONTRAST, all prenasalized segments will be voiced. On the other hand, if the ranking is reversed, voicing contrasts among prenasalized segments will be maintained. This type of approach is left open for future study.

4.3 Post-Sonorant Voicing and Pre-Sonorant Voicing

In an earlier section, we saw that there is an asymmetry in the manifestation of post-sonorant voicing. If a language has post-approximant voicing, it also has post-nasal voicing, but the reverse is not true. Languages which have only post-nasal voicing seem to be more common than languages which have both post-nasal and post-approximant voicing. However, the

same is not true for pre-sonorant voicing. In Catalan, for example, nasals in onset trigger voicing assimilation to the preceding coda; however, so do other sonorants (as well as voiced obstruents; see Section 4.3.2). Some examples are provided below.

(28) Catalan Voicing Assimilation (Mascaró 1987)

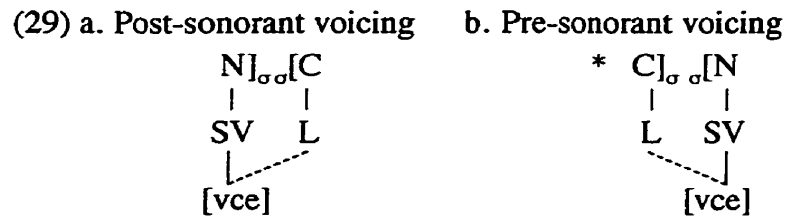
a.	me[z]os	‘months’	me[s]	‘month’
	me[s k]urt	‘short month’	me[z β]inent	‘next month’
b.	to[t]	‘all’	to[d r]ic	‘all rich persons’
c.	se[t]	‘seven’	se[d m]ans	‘seven hands’

As (28) shows, Catalan voiced obstruents undergo devoicing in coda (compare ‘me[z]os’ with ‘me[s]’). However, they remain as voiced when followed by a voiced obstruent or sonorant in onset. Catalan is representative of the following generalization: hetero-syllabic pre-sonorant voicing appears to be triggered by all sonorant consonants in all cases. As far as I know, there is no language where only nasals trigger this type of assimilation. Thus, pre-sonorant voicing and post-sonorant voicing should be analyzed differently. We will return to this shortly.

4.3.1 Feature Licensing Condition

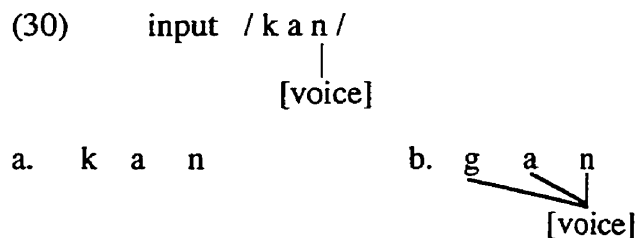
Thus far, my analysis maintains that post-sonorant voicing is a result of Licensing Cancellation. In other words, since [voice] is redundant for SV-bearing segments, it cannot be licensed by these segments and must instead be parasitically-licensed by other proper

licensors. However, by itself, this analysis does not restrict parasitic licensing to the intersyllabic sonorant + obstruent configuration. Compare the representations below.



In the post-sonorant voicing representation in (29a), [voice] cannot be licensed by the nasal because of Licensing Cancellation. Therefore, it is parasitically licensed by the following onset obstruent. However, Licensing Cancellation does not preclude the parasitic-licensing option in (29b). There is no reason why the [voice] of the onset nasal cannot be parasitically licensed by the preceding coda.

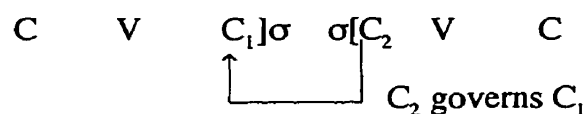
By itself, parasitic licensing also predicts that some languages should select (30b) as the optimal output for an input like /kan/, where voicing in the nasal and the vowel is licensed by the initial obstruent. However, no language has words in which all segments share [voice].



The explanation for the illformedness of (30b) proposed by Itô, Mester & Padgett (1995) appeals to the NOLINK family of constraints. However, as discussed in Section

4.2.4.1, their account suffers from both theoretical and empirical problems. To account for the absence of voicing assimilation like that exhibited in (30b), we must determine the configurations under which parasitic licensing is possible. Toward this goal, I appeal to the principle of Government from Kaye, Lowenstamm & Vergnaud (1985, 1990). In Section 2.2.3, we discussed two types of government; 1. Constituent Government, 2. Interconstituent Government. Here, we are only concerned with Interconstituent Government which is illustrated in (31): an onset governs a preceding coda.

(31) Interconstituent Government (Onset to Coda)



To restrict parasitic-licensing to occur only from coda to onset, I propose the following universal principle.

(32) Feature Licensing Condition¹⁷

A feature α which is a daughter of segment β can be licensed only by β itself, or by an interconstituent governor of β .

Since phonological government is an asymmetrical relationship, only an onset can serve as a parasitic-licenser of some property of a coda segment, not the other way around. This condition accounts for the fact that representations like (30b) are not permitted in any language. Moreover, the condition does not allow parasitic-licensing of [voice] by a segment which is not an interconstituent governor of the segment specified for [voice]. Therefore, the

Feature Licensing Condition removes pre-sonorant voicing in (29b) from the possible parasitic-licensing configurations.

Contrary to the spirit of Optimality Theory, the directional asymmetry and the restricted occurrence of voicing assimilation discussed here cannot be accounted for only by constraint reranking. I have argued that in order to account for the restrictions observed, reference to the structural relationship that holds between syllable positions is necessary.

4.3.2 Pre-Sonorant Voicing

Intersyllabic regressive voicing assimilation can be divided into two types: 1. only voiced obstruents in onset trigger voicing assimilation to the preceding coda; 2. voiced obstruents and all sonorant consonants in onset trigger voicing assimilation to the preceding coda. The first type is represented by Ukrainian which was seen in (2) and also by Warsaw Polish (Booij & Rubach 1987; Rubach 1996). The second type is represented by Catalan. In Section 4.1, I proposed that regressive assimilation is a result of the Sonority constraint. In this section, I will argue that this analysis holds for the second type as well.

I have suggested that pre-sonorant voicing be excluded from the cases of parasitic-licensing since, under the definition of interconstituent government, codas cannot be parasitic-licensors. How, then, can we capture pre-sonorant voicing? I propose that pre-sonorant voicing is a consequence of the satisfaction of SONORITY.

4.3.2.1

Sonority Constraint

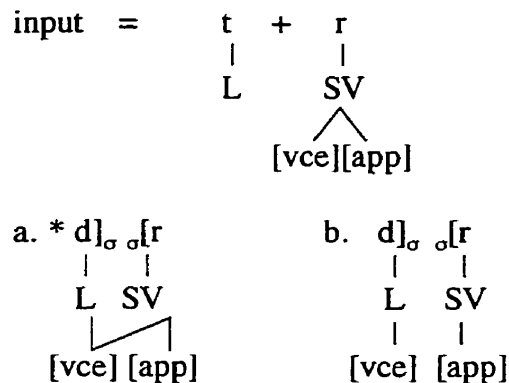
Recall from (4) in Section 4.1 that in hetero-syllabic clusters, the constraint SONORITY requires that an onset not exceed the preceeding coda in sonority value. Notice that in hetero-syllabic obstruent + sonorant sequences, SONORITY is violated. Consider the Catalan data below which are repeated from (28).

(33) Catalan Voicing Assimilation (Mascaró 1987) (= (28))

- | | | | |
|------------|---------------|--------------|---------------|
| a. me[z]os | 'months' | me[s] | 'month' |
| me[s k]urt | 'short month' | me[z β]inent | 'next month' |
| b. se[t] | 'seven' | se[d m]ans | 'seven hands' |

Recall that in Catalan, voicing assimilation is triggered both by obstruents and by sonorants. The representation below illustrates the process as triggered by an approximant.

(34) Pre-sonorant Voicing (/t + r/ → [dr])



If the /t+r/ sequence is syllabified as coda-onset, it violates sonority since /t/ has a value of zero on the sonority scale (Laryngeal node = $\emptyset$) while /r/ has a value of three (SV + [approximant] = 2+1=3). Although the presence of [voice] in the input for /r/ would appear to increase /r/'s sonority value by one, [voice] for /r/ cannot be parsed since /t/ is not in the position of an interconstituent governor. Thus, (34a) cannot be the correct output. The structure in (34a) resembles that for post-nasal voicing (cf. (29a)). However, as mentioned earlier, it is an impossible output. The illformedness of (34a) is independent of where [voice] originated in the input—in coda or onset. The problem in (34a) concerns the licensing of [voice]. As a property of the coda, Laryngeal cannot be a parasitic licenser in this configuration. Since a parasitic licenser must be an inter-constituent governor, i.e., an onset, [voice] under SV is not licensed due to the universally undominated Licensing Cancellation.

I argue instead that the structure in (34b) represents the output of pre-sonorant voicing. [voice], which is a property of /r/ in the input, is parsed by the coda [d]. Although the output in (34b) still violates SONORITY—[r] has a value of 3 while [d] has a value of 1 (therefore, there are two violations)—the degree of violation is less severe than that of the representation without [voice]-migration where the degree of violation is 3. Therefore, the migration of [voice] serves as a “partial repair” of the SONORITY violation. Although it does not fully satisfy SONORITY, it lessens the degree to which the constraint is violated. It provides a better sonority profile in the output without being unfaithful to the obstruency of the coda. This partial repair is exactly the kind of effect we would expect in Optimality Theory. A compromise between good sonority profile and faithfulness leads to partial repair.

Feature migration in (34b) should be treated differently from the unparsing of a feature, since [voice], which surfaces as a property of the coda obstruent, does have a

correspondent in the input. I suggest that feature migration violates IDENT(F), which is proposed as a member of the faithfulness constraint family together with MAX and DEP.

(35) IDENT(F) (McCarthy & Prince 1996: 264)

Let α be a segment in S_1 and β be any correspondent of α in S_2 .

If α is $[\gamma F]$, then β is $[\gamma F]$.

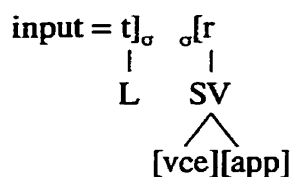
(Correspondent segments are identical in feature F).

IDENT(F) militates against changes in the featural content of a segment. Any deletion or insertion of a feature which violates MAX or DEP also violates IDENT(F). On the other hand, feature migration does not violate MAX or DEP. However, it does incur violations of IDENT(F). In (34b), there are two violations of IDENT[vce]. [voice] is lost in onset /r/ and it is added to coda /t/.

The tableaux in (36) and (37) show how the pre-sonorant voicing option is selected for a /t + r/ sequence and a /t + n/ sequence, respectively.

(36) Pre-Liquid voicing:

MAXLAR » SONORITY » IDENTVCE » IDENTSV



		MAXLAR	SONORITY	IDENTVCE	IDENTSV
	a. t] _σ σ[r L SV [app]		***	*	
	b. d] _σ σ[r L SV [vce] [app]		**	**	
	c. r] _σ σ[r ^ SV [app]	*		*	**

Although the sharing of the SV in (36c) completely satisfies SONORITY, it crucially violates MAXLAR which is the highest ranked among the three constraints. We will see shortly that this type of output is selected in Korean. When candidates (a) and (b) are compared, it can be seen that (a) violates SONORITY by three since the onset has a sonority value of 3 (SV + [approx] = 2 + 1) while the coda has a value of 0 (this candidate also violates MAXVCE which is absent from the tableau). Notice that candidate (b) does not fully satisfy SONORITY. It violates this constraint by two since the onset has a sonority value of 3 while the coda has a sonority value of 1 (L + [voice] = 0 + 1). However, (b) still emerges as optimal. Thus, the

ranking in (36) results in pre-liquid voicing. We will now consider the case of pre-nasal voicing.

(37) Pre-Nasal Voicing:
 MAXLAR » SONORITY » IDENTVCE, IDENTSV

input = $t]_{\sigma} \sigma[n$
 | |
 L SV
 |
 [vce]

		MAXLAR	SONORITY	IDENTVCE	IDENTSV
	$a.t]_{\sigma} \sigma[n$ L SV		**	*	
US	$b.d]_{\sigma} \sigma[n$ L SV [vce]		*	**	
	$c.n]_{\sigma} \sigma[n$ $\swarrow \searrow$ SV	*		*	*

Similar to the tableau in (36), SV sharing in (37c) is not selected because of its fatal violation of MAXLAR. Between candidates (a) and (b), both of which satisfy MAXLAR, the violation of SONORITY by (b) is less severe than that by (a). Therefore, (b) is selected as optimal.

Notice that the present analysis of pre-sonorant voicing predicts that if a language has pre-liquid voicing, it should also have voicing before nasals and voiced obstruents. In languages where the constraint-ranking is MAXLAR » SONORITY » IDENTVCE, a coda obstruent cannot become a sonorant to satisfy SONORITY when it is followed by a sonorant onset because losing obstruency (the Laryngeal node) will violate the higher ranked constraint, MAXLAR. Consequently, by transmitting [voice] to the preceding coda, the sonority profile is partially repaired. This ranking will select the pre-sonorant voicing option for all voiced segments, including liquids, nasals, and voiced obstruents (see tableaux (36) and (37)). Thus, under the present analysis where pre-sonorant voicing is argued to be a consequence of SONORITY, we should not find a language where only nasals trigger voicing of the preceding coda. To my knowledge, no such language is attested.

If the ranking is SONORITY » MAXLAR, we should expect coda obstruents to become sonorants by losing their Laryngeal node. Korean is a language of this type. Examples of Korean sonorantization are provided below.

(38) Korean (Cho 1990a:98)

a.	kak	+	mok	→	kaŋmok	‘stick’
b.	nap	+	nita	→	namnita	‘sprout’
c.	kat ^h	+	ni	→	kanni	‘Is it the same?’
d.	tikit	+	liil	→	tikilliil	‘the letters t and l’

It is clear from the data in (38a, b) that this process is not gemination. As all of the examples in (38) show, Korean fully satisfies SONORITY by losing coda Laryngeal and sharing SV

instead. Although the output violates MAXLAR and DEPSV, the sonority profile that results is more highly valued in this language.

If IDENT[VCE] and IDENTSV are ranked higher than the other two constraints, we should not expect pre-sonorant voicing at all. In other words, if the ranking is IDENT[VCE], IDENTSV » MAXLAR » SONORITY, the output in (37a) will be selected as optimal.

4.3.2.3 Problem of Counting

One drawback of the analysis proposed thus far lies in the formulation of SONORITY. Recall that violations of SONORITY are calculated gradiently, i.e., as the value of $\{C_2\}$ - $\{C_1\}$ increases, violation marks are accumulated. To determine the sonority value of each segment, we must add up the values assigned to relevant features and nodes. Such an interpretation of constraints is questionable since it must introduce the notion of ‘counting’ into the grammar, a mechanism which is traditionally rejected.

The problem of ‘counting’, however, seems to be inevitable when dealing with the sonority hierarchy. Most languages require a certain distance in sonority between two consonants in an onset cluster. For example, as mentioned earlier, in Attic Greek, onset clusters are restricted to “voiceless stop + n, l, r” or “voiced stop + l, r”. Clusters like “voiced stop + n” are not licit. This fact can be captured if we postulate that the sonority distance in Greek onset clusters must be at least 2. To account for such cases, researchers have proposed that values be assigned to each class of segments (Selkirk 1982 and others). There may be a way to capture sonority differences without introducing counting, but this problem is left open for further investigation.

4.4 Summary

In this chapter, we have considered different types of voicing assimilations, focusing on processes which are triggered by sonorants. I have argued that post-sonorant voicing and pre-sonorant voicing must be treated differently. I have proposed that pre-sonorant voicing is a case of feature migration and post-sonorant voicing is a case of parasitic-licensing.

In post-sonorant voicing, [voice] of SV cannot be licensed by the sonorant because of Licensing Cancellation. Thus, [voice] must be parasitically licensed by the Laryngeal node of the following obstruent. In post-sonorant voicing, assimilation is usually triggered by nasals. If approximants trigger voicing in a language, the language will also have post-nasal voicing.

The same pattern does not hold in the pre-sonorant voicing case. Assimilation is triggered by all voiced segments including voiced obstruents, nasals and approximants. Therefore, these two voicing assimilation processes must be treated differently. I have argued that only post-sonorant voicing is a case of parasitic-licensing by proposing the Feature Licensing Condition which requires the licenser of a feature to be the segment which contains it, or an interconstituent governor of that segment.

I have suggested that pre-sonorant voicing is a consequence of SONORITY which requires an onset to be less than or equal to the preceding coda in sonority. However, the appropriate formulation of this constraint is left to future research.

—NOTES—

¹ The glosses in group (a) were provided by Roumyana Slabakova (personal communication).

As in many Slavic languages, the apostrophe indicates that a consonant is ‘soft’ (palatalized), and not glottalized.

² I will return to the formulation of this constraint in Section 4.3.2.

³ Glyne Piggott suggested this constraint to me.

⁴ However, Zoque allows nasal + voiceless fricative clusters. Also, a stop can be voiceless in pre-nasal position when the stop is homorganic with the nasal.

⁵ Itô, Mester & Padgett focus on Japanese which has a process that suggests voicing underspecification for sonorants (Rendaku) as well as post-nasal voicing. I will discuss Japanese in detail in Chapter 5.

⁶ In prenasalized obstruents, voicing agreement is very common. See Section 4.2.5 for discussion.

⁷ Itô, Mester and Padgett (1995) assume that LICENSE is a constraint which is universally undominated.

⁸ However, if licensing applies only in the output level, and not in the input level, such an difference would not emerge.

⁹ The representation for a nasal in coda position will be discussed shortly.

¹⁰ I do not intend to mean that nasals (sonorants) and voiced obstruents have nothing in common. As will be discussed in Section 5.4, I will argue that both have [voice].

However, because of LC, the [voice] specification of a nasal is lost unless it is parasitically licensed, and in such a case, nasals do not behave as voiced.

¹¹ This analysis, like Itô, Mester & Padgett's (1995) analysis, cannot determine whether the output representation for post-nasal voicing involves the sharing of [voice] intersyllabically, or the passing of [voice] from the coda to the following onset with delinking of [voice] from the nasal. While the former involves fewer violations of IDENTVCE, which requires input-output identity of the feature [voice] for each segment, the latter satisfies CRISP at the cost of incurring one more violation of IDENTVCE (for the IDENT family of constraints, see Section 4.3.2.1). I have found no evidence to determine the relative ranking between IDENTVCE and CRISP in the languages which have post-nasal voicing.

¹² Only stops are voiced in this environment. When fricatives or affricates follow nasals (or nasals and laterals in Basque), they remain voiceless in both languages. In present-day Basque, morpheme-internal sequences of lateral or nasal + voiceless stop are frequently found. The rule productively affects inflectional suffixes only (Hualde 1988). Finally, 'r' does not cause voicing in Basque whereas 'l' does (e.g., /ar-tu/ → [artu] 'take' vs. /afal-tu/ → [afaldu] 'have dinner').

¹³ Itô, Mester & Padgett (1995) assume that sonorants lack a [voice] specification underlyingly. Therefore, the appearance of [voice] in the tableau in (19) is a violation of DEP_{VCE} (FILL_{VCE} in their approach since their analysis is couched within the original version of faithfulness, not within Correspondence theory). The issue of [voice] specification for sonorants will be revisited in Section 5.4.

¹⁴ The lack of [voice] spreading from vowels is not due to *COMPLEX. Rather, it is due to an independent principle which I will introduce in Section 4.3.1.

¹⁵ *COMPLEX also predicts that in languages where the feature [nasal] must be projected—languages with nasal harmony—post-nasal voicing should be unattested unless it is accompanied by post-approximant voicing. With [nasal] specified, *COMPLEX will not be able to differentiate approximants from nasals since both will have a dependent under SV in addition to [voice]. However, it seems to be the case that languages with nasal harmony either do not allow codas, do not have post-nasal voicing or do not have approximant + obstruent sequences at all.

¹⁶ Alternative representations for pre-nasalized segments have also been proposed. For example, Sagey (1986) has proposed a representation which contains only one Root node. On the other hand, representations which contain two Root nodes have been proposed by Herbert (1975, 1986), Rosenthal (1989) and others.

¹⁷ I recognize that the current formulation of Feature Licensing Condition is inadequate. In this formulation, a migrated feature (e.g., [voice] in Ukrainian) cannot be licensed. The condition in (32) is intended to limit the configuration of parasitic licensing and the better formulation will be necessary.

5 . YAMATO-JAPANESE

In the previous chapter, we discussed the problem of voicing underspecification for sonorants and how it can be resolved through Licensing Cancellation. In this chapter, I will extend the analysis which combines Licensing Cancellation and the dual-dependency of [voice] to some phenomena in Yamato-Japanese. Yamato-Japanese has attracted much attention in the recent literature because it has certain contradictory properties: *Rendaku* which requires [voice] for sonorants to be unspecified and post-nasal voicing which requires sonorants to be specified for [voice]. In this chapter, I will examine the proposed solutions that have been put forth in the recent literature to deal with this contradiction. After discussing some empirical and conceptual problems for these analyses, I will show how the proposals put forward in this thesis can capture the Yamato-Japanese facts.

5.1 Voicing Paradox in Yamato-Japanese

Rendaku in Yamato-Japanese is probably the most frequently cited phenomenon in the literature in support of voicing underspecification for sonorants. Consider the *Rendaku* data below, repeated from Chapter 2.

(1) *Rendaku*

a. ude + tokei 'wrist' 'watch'	→ udedokei 'wrist watch'
b. te + kufi 'hand' 'comb'	→ tegufi 'hand comb'
c. te + saguri 'hand' 'search'	→ tesaguri *tezaguri 'grope'
d. ai + kagi 'match' 'key'	→ aikagi *aigagi 'spare key'
e. mizu + kame 'water' 'jar'	→ mizugame 'water jar'
f. itfigo + kari 'strawberry' 'hunting'	→ itfigogari 'strawberry picking'
g. te + sawari 'hand' 'touch'	→ tezawari 'touch'

Rendaku is blocked when the second member of a compound contains a voiced obstruent. This blocking effect is analyzed by Itô & Mester (1986) to be a consequence of the Obligatory Contour Principle (OCP) (Leben 1973) which disallows two identical elements from being adjacent on a tier. Recall that, although sonorants are phonetically voiced, they pattern together with voiceless obstruents in *Rendaku*. Thus, the *Rendaku* facts support the underspecification of [voice] for sonorants.

In contrast to these data, however, Yamato-Japanese also exhibits post-nasal voicing as in (2).

(2) Verbal Inflection

- | | | |
|---------------------|-----------|----------------------------------|
| a. yom + te → yonde | 'reading' | (Itô, Mester & Padgett 1995:576) |
| b. ʃin + te → ʃinde | 'dying' | (Itô, Mester & Padgett 1995:576) |

The standard explanation for this voicing paradox appeals to rule-ordering, as was suggested by Itô & Mester (1986). Itô & Mester's solution requires that [voice] for sonorants be unspecified at the stage when *Rendaku* applies. At a later stage, before the application of post-nasal voicing, a redundancy rule, [+sonorant]→[+voice], assigns [voice] to sonorants. This account thus crucially requires that the application of *Rendaku* precede post-nasal voicing.

However, this analysis is rejected by Itô, Mester & Padgett (1995) on empirical grounds. Itô, Mester & Padgett point out that, in addition to the data on post-nasal voicing in (2), all NC (nasal-obstruent) sequences agree in voicing in Yamato-Japanese, whether they are derived or not. In (3), the nasal appears to be underlyingly sharing its voicing feature with the following obstruent.

(3) Yamato-Japanese Post-Nasal Voicing in Monomorphemic Words (Itô, Mester & Padgett 1995:575)

- | | | |
|-----------|-------------|------------|
| c. tombo | 'dragonfly' | cf. *tompo |
| d. ʃindoi | 'tired' | *ʃintoï |
| e. unzari | 'disgusted' | *unsari |
| f. kaŋgae | 'thought' | *kaŋkae |

Even more problematic for the Itô & Mester account are data of the type in (4).

spreads, while *Rendaku* is a rule which inserts the laryngeal feature [voice] onto the initial segment of the second member of a compound. This is illustrated in (5).

(5) SV Analysis

a. $\text{ſin} + \text{te} \rightarrow \text{ſin de}$ “dying”

$\swarrow \downarrow$
 SV

b. $\text{ude} + \text{tokei} \rightarrow \text{u de do ke i}$

$\downarrow$
 [v]

c. $\text{te} + \text{saguri} \rightarrow \text{te sa gu ri} \quad * \text{te za gu ri}$

$\downarrow$
 [v]

$\downarrow \downarrow$
 [v] [v]

While this analysis accounts for post-nasal voicing and *Rendaku* independently, it cannot be extended to the data in (4), where *Rendaku* and post-nasal voicing interact. See (6).

(6)

$\text{ſirooto} + \text{kangae} \rightarrow \text{ſirootokangae}$

$\swarrow \downarrow$
 SV

$\swarrow \downarrow$
 SV

$* \text{ſirootogangae}$

$\downarrow \swarrow \downarrow$
 [v] SV

If a nasal-obstruent cluster shares SV rather than [voice], the insertion of [voice] in *Rendaku* should not lead to an OCP violation. Thus, the SV-hypothesis by itself fails to account for the facts in Yamato-Japanese. We will come back to a modified version of this hypothesis in Section 5.3.

5.1.2 Itô, Mester & Padgett (1995)

Recall from Section 4.2.2 that Itô, Mester & Padgett (1995) provided an optimality-based analysis for the voicing paradox in Yamato-Japanese. They proposed a principle, Licensing Cancellation, which prohibits redundant features from being licensed. Because of this principle, sonorants cannot license [voice] by themselves; however, [voice] for sonorants can be parasitically licensed by obstruents when obstruents are adjacent to sonorants. Consequently, [voice] for sonorants can be present in the configuration of post-nasal voicing, but not in the intervocalic context as in the *Rendaku* examples in (1e-g).

Licensing Cancellation predicts that, in some languages, non-nasal sonorants will trigger voicing assimilation. To rule out this option in Yamato-Japanese, recall that Itô, Mester & Padgett formulate the NO LINK family of constraints which bans sonorant-obstruent linkages (see Section 4.2.4 for discussion). Given these assumptions, Itô, Mester & Padgett can account for *Rendaku*, post-nasal voicing, and the examples where *Rendaku* and post-nasal voicing interact, such as /ʃirooto+kangae/ in (4). However, as formulated, the NO LINK family of constraints does not take syllable structure or directionality of assimilation into consideration. As discussed in Section 4.2.5, post-sonorant voicing occurs only in specific configurations. Since in their analysis, obstruents in pre-sonorant position should also be proper licensors of [voice], we should expect the same patterns of voicing assimilation to be attested in both pre-sonorant and post-sonorant position.

More importantly, Itô, Mester & Padgett's analysis encounters an empirical problem in Yamato-Japanese as well. Coda-voiced obstruents, like nasals, spread voicing to the following onset; see (7a,b). In addition, both voiced and voiceless obstruents lose their place specification and become sonorants in coda position as seen in (7a-d).²

(7) Voicing Assimilation and Coda Sonority

coda voiced obstruent:

a. yob + te → yonde
'call' GERUND 'calling'

b. kag + te → kayde
'smell' GERUND 'smelling'

cf. yob + u → yobu
PRES. 'call'

kag + u → kagu
PRES. 'smell'

coda voiceless obstruent:

c. kak + te → kayte
'write' GERUND 'writing'

kak + u → kaku
PRES. 'write'

coda nasal:

d. kam + te → kande
'bite' GERUND 'biting'

kam + u → kamu
PRES. 'bite'

Recall that Itô, Mester & Padgett rely on the constraint SONVOI to demand the presence of [voice] in forms like (7d); they do not assume that [voice] is present in the input for sonorants. When we compare (7a,b) and (7c), we see that both voiced and voiceless obstruents surface as sonorants. However, only underlyingly voiced obstruents trigger voicing assimilation. If we rely on SONVOI, we cannot get the difference between (7b) and (7c). This suggests that voicing assimilation is triggered by the existence of a feature [voice] in the *input*, and not by the demands of a constraint like SONVOI. In the tableaux below, it is shown that Itô, Mester & Padgett's analysis cannot capture all of the data in (7). The arrows indicate the actual outputs and the hands indicate the outputs as predicted by Itô, Mester & Padgett.

(8) inputs: a. yob + te b. kag + te c. kak + te d. kam + te

|
v

|
v

		NO-GC-LINK	SONVOI	NO-LC-LINK	NO-NC-LINK
	a-1. yon + te		*		
◇	a-2. yon + de v				*
	b-1. kay + te		*		
◇	b-2. kay + de v	*			
◇	c-1. kay + te		*		
	c-2. kay + de v	*			
	d-1. kan + te		*		
◇	d-2. kan + de v				*

If SONVOI is ranked between NO-GC-LINK and NO-LC-LINK, as Itô, Mester & Padgett propose, the right outputs obtain for (8a,c) and (8d). However, this ranking will incorrectly choose (b-1) as optimal. Notice that the actual output in (8b), (b-2), has the GC linkage which is ruled out by Itô, Mester & Padgett's analysis. We could attempt to salvage their

However, by placing SONVOI above NO-GC-LINK, the ranking will now select the wrong candidate for (9c). The problem is that the constraints proposed by Itô, Mester & Padgett (1995) cannot distinguish the outputs for /kag+te/ and /kak+te/. As can be seen in the tableaux in (8) and (9), both forms in this pair always violate the same set of constraints. Thus, depending on the ranking, Itô, Mester & Padgett's strategy, that SONVOI forces [voice] insertion, will either produce outputs where both (7b) and (7c) surface with voiceless onsets, as in (8), or where they both surface with voiced onsets, as in (9). Consequently, the difference between them cannot be explained by appealing to a constraint like SONVOI. If we abandon SONVOI, we have to explain how voicing occurs in the case of (7d). To account for why voicing takes place in this form and to account for the difference between (7b) and (7c), it is necessary to assume that [voice] is specified for sonorants in the input.

In my alternative analysis, I proposed that [voice] for sonorants is specified in the input, and that MAXVCE is responsible for post-nasal voicing, as we saw in Chapter 4. The specification of [voice] in the input will be discussed further in Section 5.4.

5.2 Coda Constraints in Japanese

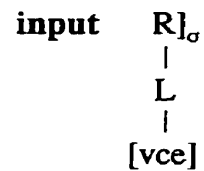
In this section, I will demonstrate that the sonorantization phenomenon can be captured with the structures and constraints which have already been proposed. Since *CODALAR prohibits Laryngeal from being in coda and SONNODE requires all segments to bear a sonority specification, satisfaction of both of these constraints will turn coda obstruents into sonorants.

Consider again the data which were introduced in (7). Notice that obstruents cannot be realized as such when they are in coda position in Japanese. In this environment, they become sonorants. In addition, Japanese does not allow place-specified segments to occupy the coda. When input labials and dorsals are syllabified as codas, they are realized as coronals (e.g., yob+te -> yonde ‘calling’) either by unparsing coda place specifications or by obtaining the place node from the following onset. These facts suggest that two coda constraints are undominated in Japanese: *CODALAR and *CODAPL.

Let us first consider coda sonorantization in detail. This process is a consequence of two constraints, one of which is *CODALAR. Recall from Chapter 3 that to be faithful to *CODALAR, some languages like German delink the Laryngeal node. Other languages like Yamato-Japanese replace coda Laryngeal with SV. The latter happens when SONNODE is highly ranked (see Section 3.2). Since SONNODE does not tolerate segments lacking both SV and Laryngeal, SV must replace Laryngeal in coda in order to be faithful to *CODALAR.

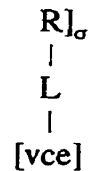
My proposal predicts the following outputs for an input coda-voiced obstruent across languages.

(10) Coda voiced obstruent



Outputs

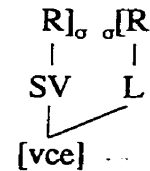
a. English



b. German



c. Japanese



When *CODALAR is ranked lower than MAXVCE (or MAXLAR) and DEPSV, the optimal output is the one in (10a) (English). When *CODALAR is ranked high, it must be respected at the expense of MAXLAR as in (10b) (German), or at the expense of DEPSV as in (10c) (Japanese). (The ranking variation will be discussed in more detail below.) When *CODALAR and MAXVCE both outrank DEP SV, the ranking forces SV insertion for a coda-voiced obstruent: since *CODALAR disallows Laryngeal from being present in the coda, only SV insertion enables [voice] to be parsed by the coda consonant. I argue that this is what happens in Yamato-Japanese.

5.3 Voicing and Sonorantization

5.3.1 Constraint-Ranking and Coda Voicing (Yamato-Japanese)

Given the proposal outlined thus far, both the voicing paradox and the sonorantization facts in Yamato-Japanese are consequences of Licensing Cancellation and *CODALAR. The tableau in (11) below shows how the optimal candidate, namely an obstruent which turns into a sonorant, is selected for an input cluster consisting of a voiced obstruent + voiceless obstruent (e.g. b+t).

(11) input= b + t
 C]σ σ[C
 | |
 R R
 | |
 L L
 |
 [vce]

	LC	*CODALAR	SONNODE	MAXVCE	DEPSV	CRISP
a. b]σ σ[t L L [vce]		*				
b. n]σ σ[d / SV L \ / [voice]					*	*
c. p]σ σ[t L			*	*		
d. n]σ σ[t / SV L				*	*	
e. n]σ σ[t / SV L \ / [vce]	*				*	

(11a) is the candidate which is identical to the input. This candidate violates *CODALAR since the representation contains a coda Laryngeal. In (11b), Laryngeal is unparsed and SV is inserted in its place; [voice] is linked to both the coda SV and the following Laryngeal. This candidate does not violate the four highest ranked constraints and

is therefore optimal. However, it violates DEP SV which militates against SV insertion. In addition, it violates CRISP which bans feature sharing over syllable boundaries. In (11c), both Laryngeal and [voice] are unparsed, so this candidate violates one of the highly-ranked constraints, MAX [voice] (as well as MAXLAR). (11d) is similar to (11b) except that in (11d), SV is inserted without satisfying MAX [voice]. Notice that by not parsing [voice], (d) does not incur a violation of CRISP. (11e) differs from (11b) in that [voice] is linked only to the coda SV and not to the onset Laryngeal. Since [voice] is redundant for SV, the representation in (11e) violates the universally undominated constraint, Licensing Cancellation. Notice that the ranking MAXVCE, SONNODE » DEPSV is crucial for Yamato-Japanese. If the ranking were reversed, (11c) would be selected as the optimal output, as is the case in German.

Before we turn to post-nasal voicing, there is one final issue that must be addressed. Candidate (11a) was the only output which we looked at where the coda maintained the Laryngeal node from the input. It is possible, however, for a coda to bear Laryngeal without violating the Coda Constraint, if Laryngeal is shared with the following onset. This configuration, though, seems to be disfavored in many languages. While voiceless geminates are relatively common, many languages do not allow voiced or aspirated geminates (e.g., Japanese, Selayarese (Mithun & Basri 1986), Korean). Therefore, a constraint which bans intersyllabic Laryngeal sharing like the one in (12) is needed.

(12) *LARSHARE⁴

Laryngeal cannot be shared across a syllable boundary.⁵

In Japanese, this constraint is undominated. Thus, voiced geminates are not allowed. If this constraint were ranked lower than *CODALAR and DEPSV, Laryngeal sharing would be chosen over sonorantization. This is the case in Ancient Greek, Dutch and Bulgarian, as was seen in Section 3.2.

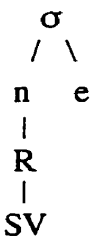
5.3.2 Post-Nasal Voicing

Let us now turn to the analysis of post-nasal voicing. Recall that in onset position, [voice] for nasals cannot be parsed because of Licensing Cancellation. In coda position, on the other hand, [voice] for nasals can be parsed parasitically by a following obstruent; the result is post-nasal voicing. This parasitic-licensing option is not available for an onset nasal since there is no proper licenser available, as defined by the Feature Licensing Condition in Section 4.3.1.

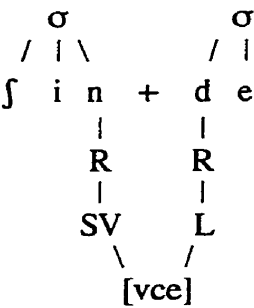
Languages like Japanese which resolve the feature licensing problem through post-nasal voicing require the ranking *CODALAR, MAXVCE » CRISP. Representations for the optimal outputs for nasals in onsets and NC sequences are in (13a) and (13b) respectively.

(13) Output Representations for nasals

a. Onset nasal

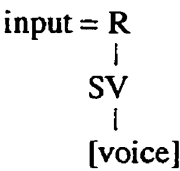


b. NC sequence



The tableaux in (14) and (15) show how various candidates are evaluated for nasals in onset and coda respectively.

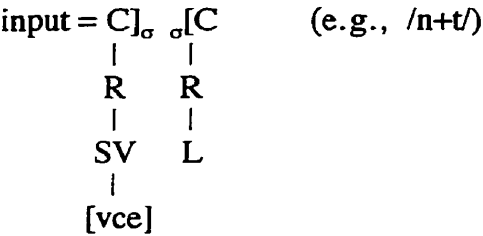
(14) Nasals in Onset



		LC	*CODALAR	MAXSV	MAXVCE
	a. R SV [vce]	*			
☞	b. R SV				*

The candidate in (14a) is identical to the input. However, it is not optimal in any language because of Licensing Cancellation being universally undominated. Instead, languages select (14b), even though it violates MAXVCE.⁶

(15) Nasals in Coda



		*CODALAR	MAXSV	MAXVCE
	a. R R SV L			*
	b. R R L		*	*
☞	c. R R SV L \ / [vce]			
	d. R R L L \ / [vce]	*	*	

In the tableau in (15), (15c) best satisfies the constraint ranking. While this candidate does not violate any of the constraints present in (15), it ultimately violates CRISP. Thus, in languages where post-nasal voicing occurs (Yamato-Japanese, Kikuyu, Zoque, etc.), CRISP

must rank lower than MAX [voice]. In languages which do not have post-nasal voicing, on the other hand, CRISP ranks higher than MAXVCE.

5.3.3 Voiced Obstruents in Coda

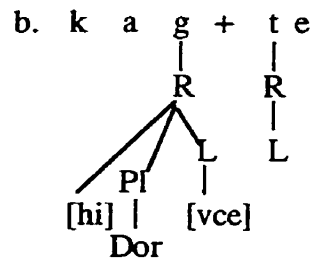
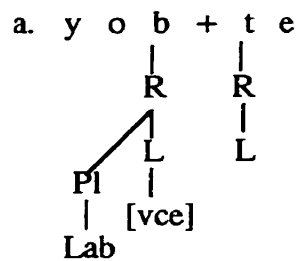
So far, we have provided analyses for post-nasal voicing and sonorantization in Yamato-Japanese. We return now to the status of voiced obstruents in coda. As pointed out in Section 5, coda voiced obstruents become sonorants when the following morpheme begins with an obstruent.

- (16) a. yob + te → [yonde]
 'call' GERUNDIVE 'calling'
 b. kag + te → [kayde]
 'smell' GERUNDIVE 'smelling'

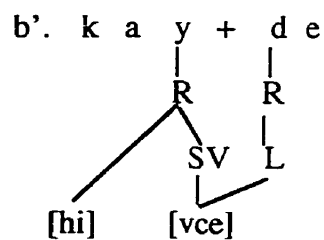
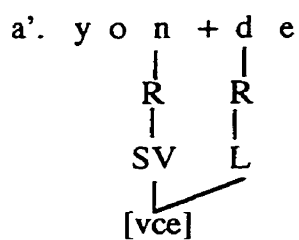
I argued earlier that SV is inserted on voiced obstruents in coda to satisfy *CODALAR in Yamato-Japanese. The proposed Feature Geometry, where [voice] is dependent on both sonority nodes, enables this feature to be parsed by the SV node in coda and to be parasitically licensed by the following Laryngeal node in onset. Output representations for input voiced obstruent + voiceless obstruent sequences are provided in (17).

(17)

Inputs

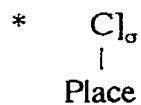


Outputs



We have yet to address the underparsing of Place features. Recall from Chapter 3 that *CODAPL prohibits the coda from licensing a Place node.

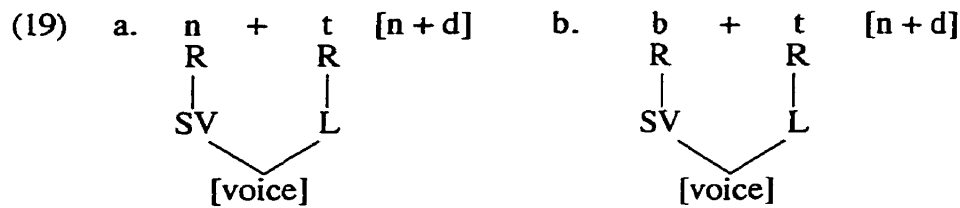
(18)*CODAPL (Itô 1986)



If *CODAPL is ranked higher than MAXPLACE in Yamato-Japanese, Place must be unparsed in coda, as illustrated in (17).⁷ In (17a'), only SV and its dependent [voice] are parsed by the coda. (17b') also contains the feature [hi]. Following Clements (1989), Wiswall (1991), Goad (1993) and others,⁸ I suggest that height features are organized under a node independent from Place. Consequently, [hi] can be retained in addition to SV and [voice], without incurring a violation of *CODAPL. Candidates in which height is parsed will thus always win over ones in which height is unparsed because the latter violate another MAX constraint. In (17a'), a sonorant stop is interpreted as a coronal nasal, along the lines of Rice (1993). In (17b'), however, the high coronal sonorant is realized as [y]. I suggest that this is because Japanese lacks the corresponding nasal [ɲ].⁹ Although Japanese allows for the homorganic velar nasal [ŋ] which would have [hi], it also has [Dorsal]_{PL}. If coda /k/ becomes [ŋ], it must share Place with the following onset because, otherwise, *CODAPL would be violated. /k/→[ŋ] would then force the onset to lose its coronal Place, thereby yielding *[kange].

5.3.4 Output Neutralization

Coda sonorantization, loss of place, and post-nasal voicing lead to neutralization of contrasts in coda-onset sequences. Compare the output for the underlying /n+t/ sequence with that for the /b+t/ sequence in (19).



Notice that the output representation for coda /b/ is identical to that for a coronal nasal. However, as we have seen, although their output representations are identical, their underlying representations must remain distinct. This difference in the input specifications is reflected in the *Rendaku* patterns in (20).

- (20) a. **No Rendaku**
 /afi + hakob + i / → afihakobi 'step' *afibakobi
 'foot' 'carry' nom.
- cf. hakob + te → hakonde
 'carry' gerundive carrying
- b. **Rendaku**
 /onna + konom + i / → onnagonomi 'women's favorite'
 'woman' 'like' nom.
- cf. konom + te → kononde
 'like' gerundive liking

In the compounding examples above, the voiced obstruent /b/ and nasal /m/ are syllabified as onsets. Consequently, the coda constraints are not applicable. Notice that the /b/ in (20a) blocks the application of *Rendaku* while the /m/ in (20b) does not. This

difference clearly shows that the representations in (19a, b) have different underlying sources, although the difference is neutralized on the surface before obstruents.

5.4 [voice] Specification for Sonorants—Against Lexicon Optimization

So far, I have demonstrated how coda-sonorantization and post-nasal voicing in Yamato-Japanese can be accounted for through the interaction of coda constraints and Licensing Cancellation. However, contrary to the position taken by Itô, Mester & Padgett (1995), I have assumed that [voice] is present underlyingly for sonorants. In this section, I will show that this assumption is needed to account for the Japanese facts.

Itô, Mester & Padgett (1995) have chosen not to specify [voice] underlyingly, since both [voice]-specified, and non-specified inputs yield the same optimal output for the data they consider. In such a case, Lexicon Optimization (Prince & Smolensky 1993:192) chooses the input which is most harmonic to the output representation.

(21) Lexicon Optimization (Prince & Smolensky 1993:192)

Suppose that several different inputs $I_1, I_2, \dots, I_N$ when parsed by a grammar G lead to corresponding outputs $O_1, O_2, \dots, O_N$, all of which are realized as phonetic form Φ — these inputs are all *phonetically equivalent* with respect to G . Now one of these outputs must be the most harmonic, by virtue of incurring the least significant violation marks: suppose this optimal one is labelled O_k . Then the learner should choose, as the underlying form for Φ , the input I_k .

One of the tenets of Optimality Theory holds that, regardless of input specification, the constraint ranking will select the correct optimal output (the Richness of the Base hypothesis). Since there are many possible input specifications which will lead to the correct input-output pairing. Lexicon Optimization in (21) fulfills this role. Among inputs, Lexicon Optimization selects as most harmonic the one which involves the fewest violations of highly ranked constraints. This is illustrated in Itô, Mester & Padgett's "tableau des tableaux" in (22), where outputs for two different inputs for /maki/ 'firewood' are evaluated.

(22) Tableau des tableaux¹⁰

	input	output	LICENSE	SONVOI	PARSE VCE	FILLVCE
	a. maki v	 maki <v>		*	*	
⇒	b. maki	 maki		*		

It can be seen that the input which lacks [voice] leads to the correct output with the fewest number of constraint violations. Consistent with Lexicon Optimization, then, Itô, Mester & Padgett (1995) argue that the input in (22b) is the right one, as it is more harmonic with the output than the input in (22a).

Recall from Section 5.1.2, however, that there are some data from Japanese which suggest the reverse, that sonorants are underlyingly specified for [voice]. The presence of [voice] is crucial in accounting for the difference between /kam+te/→[kande] and

/kag+te/→[kayde] versus /kak+te/→[kayte]: voicing assimilation is triggered only by underlying voiced obstruents and nasals. As we saw in (7), coda /g/ becomes [y] and spreads its voicing to the following onset. If [voice] is not specified for /m/ in /kam+te/ and SONVOI is responsible for post-nasal voicing, then SONVOI should also ‘insert’ [voice] on [y] in the /kak+te/→[kayte] case and incorrectly yield [kayde]. If we abandon SONVOI, something else must be responsible for voicing in the case of /kam+te/→[kande]. Consider again the data in (23).

(23)

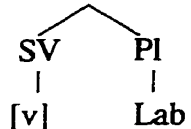
- a. kam + te → [kande]
“chew” GERUNDIVE
- b. kag + te → [kayde]
“smell” GERUNDIVE
- c. kak + te → [kayte]
“write” GERUNDIVE

Since in Optimality Theory, it is output representations which are evaluated and there is no serial derivation, the difference between the presence of voicing in (23b) and its absence in (23c) must be attributed to the input. There cannot be a constraint like SONVOI which demands that sonorants bear [voice] in the output. Consequently, to obtain post-nasal voicing in (23a), the root-final /m/ must have [voice] in the input. I suggest that the difference between the surface forms in (23b) and (23c) is due to the location of the constraint DEP_{VCE} in the ranking. Provided that sonorants have [voice] underlyingly, the ranking MAX_{VCE}, DEP_{VCE} » CRISP guarantees that the correct surface forms will be selected. See (24).¹¹

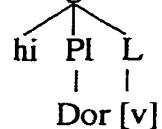
(24) Input Sonorants have [voice]:

inputs

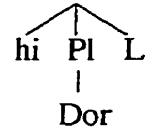
a. kam + te



b. kag + te



c. kak + te



	candidates	MAXVCE	DEPVCE	CRISP
	a.1 ka n de SV L \ / [vce]			*
	a.2 kan te SV L	*		
	b.1 kay de hi SV L \ / [vce]			*
	b.2 kay te hi SVL	*		
	c.1 kay de hi SV L \ / [vce]		*	*
	c.2 kay te hi SV L			

To summarize, we have seen that if sonorants have [voice] in the input, the right surface forms obtain. If, on the other hand, sonorants are not specified for [voice] underlyingly, as Itô, Mester & Padgett (1995) assume, we cannot account for why the SONVOI constraint

does not provide [voice] for [kayte] in (24c) while it does for [kande] and [kayde] in (24a,b); all three codas are sonorants in the output.

Although under Itô, Mester & Padgett's (1995) analysis, Lexicon Optimization selects the input where sonorants do not bear [voice], when all the facts are considered, it becomes clear that inputs where [voice] is specified for sonorants and inputs where [voice] is not specified do not select the same optimal outputs. Since only the former leads to the right surface representation, [voice] must be specified underlyingly for sonorants, contrary to their assumption.

My analysis challenges the Richness of the Base hypothesis. Although this hypothesis holds that the degree of input specification should not matter, we have seen that the presence of [voice] for input sonorants is crucial in accounting for the facts discussed here. In order to provide a comprehensive account for voicing assimilation in Japanese, we have to attribute the fact that underlying sonorants pattern together with voiced obstruents to the [voice] specification in the input. If we opt instead for the approach that there is [voice] underspecification plus a SONVOI constraint, [voice] will be inserted in [kayte] (= /kak+te/) as well as in [kande] (= /kam+te/) since, in both cases, codas surface as sonorants. Thus, the combination of Richness of the Base and Lexicon Optimization cannot be maintained¹² when we consider the full range of facts from Yamato-Japanese.

5.5 Crosslinguistic Consequences of NoCoda: Laryngeal

We have seen that Japanese chooses sonorantization in order to be faithful to *CODALAR. As discussed in Section 5.3, the ranking *CODALAR, SONNODE, MAXVCE, *LARSHARE »

DEPSV, CRISP selects sonorantization as the optimal output for coda voiced obstruents (as well as for voiceless obstruents). As expected, reranking of these constraints yields different optimal outputs. Some examples of different rankings are given in (25) below.

(25) Constraint Rankings and Possible Outputs

- a. MAXVCE, DEPSV, SONNODE, *LARSHARE, CRISP » *CODALAR
 — English
- b. *CODALAR, DEPSV, SONNODE, MAXVCE » CRISP, *LARSHARE
 — Dutch
- c. *CODALAR, DEPSV, *LARSHARE, CRISP » SONNODE, MAXVCE
 — German
- d. CRISP, *CODALAR, SONNODE, *LARSHARE » DEPSV, MAXVCE
 — Hausa
- e. *CODALAR, SONNODE, MAXVCE, *LARSHARE » DEPSV, CRISP
 — Japanese

The ranking in (25a) selects an English-type output where Laryngeal in coda is maintained on the surface. (25b) selects a Dutch-type output where Laryngeal is shared with the following obstruent. The ranking in (25c) chooses a German-type output where coda Laryngeal is neutralized and the coda bears no Sonority node on the surface. The ranking in (25d) selects a Hausa-type output where coda Laryngeal is neutralized and SV is inserted in its place. In this type of language, coda obstruents become sonorants; however, unlike in a Japanese-type language in (25e), the [voice] specification in coda is not shared with the following onset due to the highly ranked CRISP constraint.

In a theory which allows free reranking of constraints, as the number of constraints increases, the number of possible rankings does as well. However, allowing many possible

constraint rankings does not always lead to a huge number of grammars with distinct optimal outputs. For example, free reranking of the 6 constraints in (25) yields 720 ranking possibilities; however, grammars with distinct optimal outputs are limited to the five types listed in (25). Out of 720 rankings, 120 yield the same output as (25a). Similarly, for each of (25b-e), there are 48 rankings that produce the same output.¹³

Unparsing of [voice] only is also a possible output generated by GEN; however, it violates MAXVCE, in addition to *CODALAR and such a representation will never be selected over a German-type representation where both [voice] and Laryngeal are unparsed. Because of Licensing Cancellation, representations where [voice] is exclusively parsed by SV will never be selected as optimal. Thus, in order to satisfy MAXVCE and *CODALAR at the same time, there are only two possibilities; [voice] is parsed by Laryngeal which is shared with the following onset (Dutch), or SV replaces coda Laryngeal and [voice] is linked to both coda SV and onset Laryngeal (Japanese). Thus, the combination of the constraints proposed and feature geometry excludes some possibilities with the result that reranking of the 6 constraints yields only the 5 attested types of languages for coda voiced obstruents.

5.6 Summary

In this chapter, we have discussed various phenomena in Yamato-Japanese. First, I have demonstrated that the proposed geometry where [voice] is dependent on SV and Laryngeal, together with Itô, Mester & Padgett's (1995) Licensing Cancellation, solves the problem of [voice] specification in sonorants. I then discussed sonority restrictions and voicing assimilation in Yamato-Japanese. Although these appear to be unrelated, I have shown that the geometry and *CODALAR, together with other independently motivated constraints, can

account for the interaction of post-nasal voicing and coda sonorantization as related phenomena.

Related to the Yamato-Japanese phenomena, I provided support for the assumption that [voice] must be specified for sonorants in the input. This assumption challenges one of the main tenets of Optimality Theory, the Richness of the Base Hypothesis. In Optimality Theory, it is argued that whatever the input specification is, constraint ranking will select the right output—Richness of the Base. Among the set of possible input specifications, the one which is most harmonic to the optimal output is selected as the actual input—Lexicon Optimization. However, the Yamato-Japanese facts from sonorantization and voicing assimilation show that [voice] specification in the input for sonorants is crucial in accounting for the voicing assimilation patterns observed. Therefore, the Richness of the Base Hypothesis does not hold.

Finally, I showed that reranking of the constraints which accounted for Yamato-Japanese sonorantization can also account for other attested types of languages: English where Laryngeal is allowed in coda, Dutch where the coda receives Laryngeal from the following onset, German where the coda loses Laryngeal, and Hausa where the coda obstruent becomes sonorant.

—NOTES—

¹ Rice assumes that SV is a voicing feature for sonorants. She does not assume that [voice] can be dependent on SV.

² In Yamato-Japanese, most suffixes which attach to verbs begin either with /t/ or with a vowel. /t/-initial suffixes such as 'ta' (past tense) and 'tari' (representative) pattern with 'te' (gerundive) illustrated here. In addition to the obstruents in (7), Yamato-Japanese also has /t/ as a stem-final obstruent. In the case of stem-final /t/, the output of gerundive suffixation is a geminate, because all the features of stem-final /t/ can be licensed by the onset (e.g., hanat 'release' + te GERUND → hanatte 'releasing').

³Regarding NO-LC-LINK, Japanese appears to have /r/ in coda underlyingly; however, the behavior of this segment makes its status unclear. When stem-final /r/ is followed by /t/, the /t/ undergoes gemination (e.g., /kar/ 'borrow' + /te/ GERUND → [katte]; cf. /kar/ 'borrow' + /u/ PRESENT → [karu] 'borrow'). From this fact, /r/ seems to have no melodic content. In support of [r] underspecification, we find that when a vowel-initial suffix is added to a CV verbal stem, /r/ is inserted to satisfy the ONSET constraint in Japanese (a syllable must have an onset) (see Sakai 1994). I leave this problem to future research.

⁴ This constraint may be subsumed under CRISP. However, languages where Laryngeal sharing across syllable boundaries is prohibited sometimes allow Place or/and [voice] sharing across syllable boundaries (Japanese, Korean, Selayarese). In addition, sharing of certain nodes and features across syllables seems to be disfavored while sharing of others is common. For example, Place sharing seems to be far more common than

Laryngeal sharing. We cannot account for such an observation through constraint interaction if we subsume *LARSHARE under CRISP.

⁵ This constraint not only bans heterosyllabic obstruent clusters where Laryngeal is shared, but also geminates which bear Laryngeal. This constraint is designed to capture the crosslinguistic preference for voiceless+voiceless obstruent clusters over voiced+voiced ones. However, there are also languages which require obstruents over syllable boundaries to share laryngeal properties. As argued in Chapters 3 and 4, Laryngeal sharing is one way to escape from a *CODALAR violation since coda Laryngeal is licensed by the following onset. Such a scenario apparently violates the *LARSHARE constraint proposed here.

⁶ There is a way that [voice] can be parsed for an input nasal in onset. By inserting Laryngeal, sonorants can parse [voice] by turning into prenasalized voiced obstruents (cf. (24b) in Chapter 4). Outputs of this kind are attested. For example, in Amahhuaca (Osborn 1948), all intervocalic nasals become prenasalized voiced obstruents. Similar phenomena are attested in Sirionó (Firestone 1965) and Jukun (Welmers 1973).

The proposed analysis, however, also predicts the existence of languages where onset sonorants becomes full obstruents rather than prenasalized obstruents in order to parse [voice]. Whether or not such languages exist is unknown to me. If such languages prove to be uncommon or unattested, it might be due to the hypothesis that faithfulness in onset position must be respected over that of in other positions.

⁷One might ask why Place cannot be licensed by the onset. Crosslinguistically, onset specifications are usually maintained while coda specifications are often neutralized. Place assimilation from onset to coda is attested in many languages, but place assimilation from coda to onset is extremely rare. One way to capture this would be with a constraint which

holds that onset features cannot be changed. In coda to onset voicing assimilation (including Japanese post-nasal voicing), this constraint is violated, so MAXVCE and *CODALAR must be ranked higher than this constraint.

⁸ Although neither Clements nor Goad argues that velars have a height specification, outside of voicing, [g] and [y] only have the feature [high] in common. The same type of velar-high sonorant alternation is observed in Romansh, Cypriot Greek, and diachronically in English. Romansh has a process of glide hardening: /veyr/ → [vékr] ‘true’ (Cho & Inkelas 1993:6). In Cypriot Greek, the palatal glide /y/ becomes [kʲ] when it follows a consonant (Newton 1972, Kaisse 1992). In Old English, [u] > [i] > [g] (Jones 1978).

⁹ It is thus clear that a language’s inventory controls the interpretation of representations; cf. Structure Preservation (Kiparsky 1985).

¹⁰ Recall from Chapter 4 that Itô, Mester & Padgett (1995) use the original Optimality version of faithfulness constraints where PARSE is roughly equivalent to MAX and FILL is equivalent to DEP in Correspondence Theory. Although they collapse the PARSE and FILL family of constraints into FAITH, I have listed each member of this family independently.

¹¹ In the tableaux in (24), two different surface representations for [y] are provided; (b.1) with [voice] and (c.2) without [voice]. This difference is tied to their different inputs: (b) has [voice] underlyingly (voiced obstruent) while (c) does not (voiceless obstruent).

¹² Idsardi (1997) also argues against Lexicon Optimization. His evidence comes from various fields including acquisition and experimental psycholinguistics.

¹³ Other examples include;

I. English-type outputs

DEPSV, MAXVCE, SONNODE, CRISP » *CODALAR » *LARSHARE (24 rankings)

(DEPSV » MAXVCE » SONNODE » CRISP » *CODALAR » *LARSHARE)
(MAXVCE » DEPSV » SONNODE » CRISP » *CODALAR » *LARSHARE)
(DEPSV » SONNODE » MAXVCE » CRISP » *CODALAR » *LARSHARE)
(SONNODE » DEPSV » MAXVCE » CRISP » *CODALAR » *LARSHARE)
(MAXVCE » SONNODE » DEPSV » CRISP » *CODALAR » *LARSHARE)
(SONNODE » MAXVCE » DEPSV » CRISP » *CODALAR » *LARSHARE)

(CRISP » MAXVCE » SONNODE » DEPSV » *CODALAR » *LARSHARE)
(MAXVCE » CRISP » SONNODE » DEPSV » *CODALAR » *LARSHARE)
(CRISP » SONNODE » MAXVCE » DEPSV » *CODALAR » *LARSHARE)
(SONNODE » CRISP » MAXVCE » DEPSV » *CODALAR » *LARSHARE)
(MAXVCE » SONNODE » CRISP » DEPSV » *CODALAR » *LARSHARE)
(SONNODE » MAXVCE » CRISP » DEPSV » *CODALAR » *LARSHARE)

(DEPSV » CRISP » SONNODE » MAXVCE » *CODALAR » *LARSHARE)
(CRISP » DEPSV » SONNODE » MAXVCE » *CODALAR » *LARSHARE)
(DEPSV » SONNODE » CRISP » MAXVCE » *CODALAR » *LARSHARE)
(SONNODE » DEPSV » CRISP » MAXVCE » *CODALAR » *LARSHARE)
(CRISP » SONNODE » DEPSV » MAXVCE » *CODALAR » *LARSHARE)
(SONNODE » CRISP » DEPSV » MAXVCE » *CODALAR » *LARSHARE)

(DEPSV » MAXVCE » CRISP » SONNODE » *CODALAR » *LARSHARE)
(MAXVCE » DEPSV » CRISP » SONNODE » *CODALAR » *LARSHARE)
(DEPSV » CRISP » MAXVCE » SONNODE » *CODALAR » *LARSHARE)
(CRISP » DEPSV » MAXVCE » SONNODE » *CODALAR » *LARSHARE)
(MAXVCE » CRISP » DEPSV » SONNODE » *CODALAR » *LARSHARE)
(CRISP » MAXVCE » DEPSV » SONNODE » *CODALAR » *LARSHARE)

DEPSV, MAXVCE, SONNODE, *LARSHARE » *CODALAR » CRISP (24 rankings)

CRISP, MAXVCE, SONNODE, *LARSHARE » *CODALAR » DEPSV (24 rankings)

DEPSV, CRISP, SONNODE, *LARSHARE » *CODALAR » MAXVCE (24 rankings)

DEPSV, CRISP, MAXVCE, *LARSHARE » *CODALAR » SONNODE (24 rankings)

II. Dutch-type outputs

MAXVCE, SONNODE, *CODALAR » CRISP » DEPSV » *LARSHARE (6 rankings)

(MAXVCE » SONNODE » *CODALAR » CRISP » DEPSV » *LARSHARE)
(SONNODE » MAXVCE » *CODALAR » CRISP » DEPSV » *LARSHARE)
(*CODALAR » SONNODE » MAXVCE » CRISP » DEPSV » *LARSHARE)
(SONNODE » *CODALAR » MAXVCE » CRISP » DEPSV » *LARSHARE)
(MAXVCE » *CODALAR » SONNODE » CRISP » DEPSV » *LARSHARE)
(*CODALAR » MAXVCE » SONNODE » CRISP » DEPSV » *LARSHARE)

DEPSV, SONNODE, *CODALAR » CRISP » MAXVCE » *LARSHARE (6 rankings)
MAXVCE, SONNODE, *CODALAR » CRISP » SONNODE » *LARSHARE (6 rankings)
DEPSV, SONNODE, *CODALAR » *LARSHARE » MAXVCE » CRISP (6 rankings)
MAXVCE, DEPSV, *CODALAR » *LARSHARE » SONNODE » CRISP (6 rankings)
MAXVCE, SONNODE, *CODALAR » *LARSHARE » DEPSV » CRISP (6 rankings)
DEPSV, SONNODE, *CODALAR » CRISP » *LARSHARE » MAXVCE (6 rankings)
DEPSV, SONNODE, *CODALAR » *LARSHARE » CRISP » MAXVCE (6 rankings)
DEPSV, MAXVCE, *CODALAR » CRISP » *LARSHARE » SONNODE (6 rankings)
DepSV, MaxVce, *CodaLar » *LarShare » Crisp » SonNode (6 rankings)

III. German-type outputs

*CODALAR, CRISP, DEPSV » SONNODE » MAXVCE » *LARSHARE (6 rankings)

(*CODALAR » CRISP » DEPSV » SONNODE » MAXVCE » *LARSHARE)
(CRISP » *CODALAR » DEPSV » SONNODE » MAXVCE » *LARSHARE)
(*CODALAR » DEPSV » CRISP » SONNODE » MAXVCE » *LARSHARE)
(DEPSV » *CODALAR » CRISP » SONNODE » MAXVCE » *LARSHARE)
(*CODALAR » CRISP » DEPSV » SONNODE » MAXVCE » *LARSHARE)
(CRISP » *CODALAR » DEPSV » SONNODE » MAXVCE » *LARSHARE)

*CODALAR, CRISP, DEPSV » MAXVCE » SONNODE » *LARSHARE (6 rankings)
*CODALAR, *LARSHARE, DEPSV » SONNODE » MAXVCE » CRISP (6 rankings)
*CODALAR, *LARSHARE, DEPSV » MAXVCE » SONNODE » CRISP (6 rankings)
*CODALAR, CRISP, DEPSV » SONNODE » *LARSHARE » MAXVCE (6 rankings)
*CODALAR, *LARSHARE, DEPSV » SONNODE » CRISP » MAXVCE (6 rankings)
*CODALAR, *LARSHARE, DEPSV » MAXVCE » CRISP » SONNODE (6 rankings)
*CODALAR, CRISP, DEPSV » MAXVCE » *LARSHARE » SONNODE (6 rankings)
*CODALAR, *LARSHARE, CRISP » MAXVCE » DEPSV » SONNODE (6 rankings)

IV. Hausa-type outputs

*CODALAR, CRISP, SONNODE» MAXVCE» DEPSV » *LARSHARE (6 rankings)

(*CODALAR» CRISP» SONNODE» MAXVCE» DEPSV » *LARSHARE)
(CRISP » *CODALAR» SONNODE» MAXVCE» DEPSV » *LARSHARE)
(*CODALAR» SONNODE» CRISP» MAXVCE» DEPSV » *LARSHARE)
(SONNODE» *CODALAR» CRISP» MAXVCE» DEPSV » *LARSHARE)
(CRISP » SONNODE» *CODALAR» MAXVCE» DEPSV » *LARSHARE)
(SONNODE» CRISP » *CODALAR» MAXVCE» DEPSV » *LARSHARE)

*CODALAR, CRISP, SONNODE» DEPSV » MAXVCE » *LARSHARE (6 rankings)

*CODALAR, CRISP, SONNODE» MAXVCE» *LARSHARE» DEPSV (6 rankings)

*CODALAR, CRISP, *LARSHARE» MAXVCE» SONNODE» DEPSV (6 rankings)

*CODALAR, CRISP, SONNODE» DEPSV » *LARSHARE» MAXVCE (6 rankings)

*CODALAR, *LARSHARE, SONNODE» DEPSV » CRISP » MAXVCE (6 rankings)

V. Japanese-type outputs

*CODALAR, *LARSHARE, SONNODE» DEPSV » MAXVCE» CRISP (6 rankings)

(*CODALAR» *LARSHARE» SONNODE» DEPSV » MAXVCE» CRISP)
(*LARSHARE» *CODALAR » SONNODE» DEPSV » MAXVCE» CRISP)
(SONNODE» *LARSHARE» *CODALAR» DEPSV » MAXVCE» CRISP)
(*LARSHARE» SONNODE» *CODALAR» DEPSV » MAXVCE» CRISP)
(SONNODE» *CODALAR» *LARSHARE» DEPSV » MAXVCE» CRISP)
(*CODALAR» SONNODE» *LARSHARE» DEPSV » MAXVCE» CRISP)

*CODALAR, *LARSHARE, MAXVCE» DEPSV » SONNODE» CRISP (6 rankings)

*CODALAR, *MAXVCE, SONNODE» CRISP » *LARSHARE» DEPSV (6 rankings)

*CODALAR, *LARSHARE, MAXVCE» CRISP » SONNODE» DEPSV (6 rankings)

*CODALAR, *LARSHARE, MAXVCE» DEPSV » CRISP » SONNODE (6 rankings)

*CODALAR, *LARSHARE, MAXVCE» CRISP » DEPSV » SONNODE (6 rankings)

REFERENCES

- Anderson, J. and C. Ewen. 1987. *Principles of Dependency Phonology*. Cambridge: Cambridge University Press.
- Archangeli, D. and D. Pulleyblank. 1989. Yoruba Vowel Harmony. *Linguistic Inquiry*. 20:173-219.
- Archangeli, D. and D. Pulleyblank. 1994. *Grounded Phonology*. Cambridge, MA: MIT Press.
- Armstrong, L. E. 1967. *The Phonetic and Tonal Structure of Kikuyu*. Cambridge: Cambridge University Press.
- Avery, P. and Rice K. 1989. Segment Structure and Coronal Underspecification. *Phonology* 6:179-200.
- Barker, M. A. R. 1963. *Klamath Dictionary*. Berkeley: University of California Press.
- Barker, M. A. R. 1964. *Klamath Grammar*. Berkeley: University of California Press.
- Bendor-Samuel, J. T. 1960. Some Problems of Segmentation in the Phonological Analysis of Tereno. *Word*. 16: 348-355.
- Bethin, C. Y. 1992. Intonation and Gemination in Ukrainian. *Slavic and East European Journal*. 36: 275-301.
- Blevins, J. 1995. The Syllable in Phonological Theory. In J. Goldsmith (Ed.), *A Handbook of Phonological Theory*. 206-244. Cambridge: Blackwell.
- Booij, G. E. and J. Rubach. 1987. Postcyclic Versus Postlexical Rules in Lexical Phonology. *Linguistic Inquiry*. 18: 1-44.

Broselow, E. 1995. Skeletal Positions and Moras. In J. Goldsmith (Ed.), *The Handbook of Phonological Theory*. 175-205. Cambridge, MA: Basil Blackwell.

Carnie, A. 1994. Whence Sonority? Evidence from Epenthesis in Modern Irish. In A. Carnie, H. Harley, and T. Bures (Eds.), *MIT Working Papers in Linguistics Vol. 21: Papers on Phonology and Morphology*. 81-108.

Causley, T. 1996 . Featural Correspondence and Identity: the Athapaskan Case. Paper presented at NELS 27.

Cho, Y. 1990a. *Parameters of Consonantal Assimilation*. Ph. D. Dissertation, Stanford University.

Cho, Y. 1990b. A Typology of Voicing Assimilation. *Proceedings of WCCFL IX*. 141-155.

Cho, Y. 1991. Voicing is not Relevant for Sonority. *BLS 17*. 69-80.

Cho, Y. and S. Inkelas. 1993. Major Class Alternations. In E. Duncan, D. Farkas, and P. Spaelti (Eds.), *Proceedings of WCCFL XII*. 3-18: CSLI.

Chomsky, N. 1957. *Syntactic Structures*. The Hague: Mouton.

Chomsky, N. 1965. *Aspects of the Theory of Syntax*. Cambridge: MIT Press.

Chomsky, N. 1981. *Lectures on Government and Binding*. Dordrecht: Foris.

Chomsky, N. 1986. *Barriers*. Cambridge: MIT Press.

Chomsky, N. 1989. Some Notes on Economy of Derivation and Representation. In I. Laka and A. Mahajan (Eds.), *MIT Working Papers in Linguistics 10* 43-47. Cambridge: MIT Department of Linguistics and Philosophy.

Chomsky, N. and M. Halle. 1968. *The Sound Pattern of English*. New York: Harper and Row.

Clements, G. N. 1985. The Geometry of Phonological Features. *Phonology*. 2: 225-252.

Clements, G. N. 1990. The Role of the Sonority Cycle in Core Syllabification. In J. Kingston and M. Beckman (Eds.), *Papers in Laboratory Phonology 1: Between the Grammar and Physics of Speech*. 283-333. Cambridge: Cambridge University Press.

Clements, G. N. and E. Hume. 1995. The Internal Organization of Speech Sounds. In J. Goldsmith (Ed.), *A Handbook of Phonological Theory*. 244-305. Cambridge: Blackwell.

Cole, J. and C. Kisseberth. 1994 . An Optimal Domain Theory of Harmony. Paper presented at the FLSM V.

Davy, J. J. M and D. Nurse. 1982. Synchronic Version of Dahl's Law: The Multiple Applications of a Phonological Dissimilation Law. *Journal of African Languages and Linguistics*. 4:157-195.

Firestone, H. L. 1965. *Description and Classification of Sirionó*. The Hague: Mouton.

Fleming, I. and R. K. Dennis. 1977. Tol (Jicaque) Phonology. *International Journal of American Linguistics*. 43: 121-127.

Giegrich, H. 1985. *Metrical Phonology and Phonological Structure*. Cambridge: Cambridge University Press.

Goad, H. 1993. On the Configuration of Height Features. Ph. D. Dissertation, University of Southern California.

Goad, H. 1996. Consonant Harmony in Child Language: Evidence against Coronal Underspecification. In B. Bernhardt, J. Gilbert, and D. Ingram (Eds.), *Proceedings of the UBC International Conference on Phonological Acquisition*. 187-200. Somerville, MA: Cascadia Press.

Goad, H. in press. Consonant Harmony in Child Language: An Optimality-Theoretic Account. In S.-J. Hannahs and M. Young-Scholten (Eds.), *Focus on Phonological Acquisition*. Amsterdam: John Benjamins.

Goldsmith, J. 1993. Harmonic Phonology. In J. Goldsmith (Ed.), *The Last Phonological Rule: Reflections on Constraints and Derivations*. 21-60. The University of Chicago Press.

Greenberg, J. H. 1978. Some Generalizations Concerning Initial and Final Consonant Clusters. In J. H. Greenberg (Ed.), *Universals of Human Language*. Vol. 2, 243-279. Stanford: Stanford University Press.

Groves, T. R., G. W. Groves and R. Jacobs. 1985. *Kiribatese: An Outline Description*. Canberra: The Australian National University.

Halle, M. 1959. *The Sound Pattern of Russian*. The Hague: Mouton.

Halle, M. and K. Stevens. 1971. A Note on Laryngeal Features. Quarterly Progress Report. 101: 198-212. Cambridge: Research Laboratory of Electronics, MIT.

Haraguchi, S. 1984. Some Tonal and Segmental Effects of Vowel Height in Japanese. In M. Aronoff and R. T. Oehrle (Eds.), *Language Sound Structure*. 143-156. Cambridge: MIT Press.

Harris, J. 1990. Segmental Complexity and Phonological Government. *Phonology*. 7: 255-300.

Harris, J. 1994. *English Sound Structure*. Oxford: Blackwell Publishers.

Harris, J. W. 1986. Autosegmental Phonology and Liquid Assimilation in Havana Spanish. In L. D. King and C. A. Maley (Eds.), *Selected Papers from the XIIIth Linguistic Symposium on Romance Languages (1983)*. 127-148. Amsterdam/Philadelphia: John Benjamins.

Hayes, B. 1986. Inalterability in CV Phonology. *Language*. 62: 321-351.

Hayes, B. 1989. Compensatory Lengthening in Moraic Phonology. *Linguistic Inquiry*. 20: 253-306.

Herbert, R. K. 1975. Reanalyzing Prenasalized Consonants. *Studies in African Linguistics*. 6: 105-123.

Herbert, R. K. 1986. *Language Universals, Markedness Theory, and Natural Phonetic Processes*. Amsterdam: Mouton de Gruyter.

Hooper, J. B. 1976. *An Introduction to Natural Generative Phonology*. New York: Academic Press.

Hualde, J. I. 1988. *A Lexical Phonology of Basque*. Ph. D. Dissertation, University of Southern California.

Huffman, M. K. 1993. Phonetic Patterns of Nasalization and Implications for Feature Specification. In M. K. Huffman and R. A. Krakow (Eds.), *Phonetics and Phonology 5: Nasals, Nasalization, and the Velum*. 303-327. San Diego: Academic Press.

Humbert, H. 1996. On the Asymmetrical Nature of Nasal Obstruent Relations. Paper presented at NELS 27, McGill University.

Humesky, A. 1980. *Modern Ukrainian*. Canadian Institute of Ukrainian Studies.

Hyman, L. 1984. On the Weightlessness of Syllable Onsets. In C. Brugman and M. Macaulay (Eds.), *BLS 10*.

Hyman, L. 1985. *A Theory of Phonological Weight*. Dordrecht: Foris.

Idsardi, W. 1997. Against Lexicon Optimization. Paper presented at MOT Phonology Workshop.

Itô, J. 1986. *Syllable Theory in Prosodic Phonology*. Ph.D. Dissertation, University of Massachusetts, Amherst.

Itô, J. and A. Mester. 1986. The Phonology of Voicing in Japanese. *Linguistic Inquiry*. 17: 45-73.

Itô, J. and A. Mester. 1994. Reflections on CodaCond and Alignment. In J. Merchant, J. Padgett, and R. Walker (Eds.), *Phonology at Santa Cruz III*. 27-46.

Itô, J., A. Mester and J. Padgett. 1995. NC: Licensing and Underspecification in Optimality Theory. *Linguistic Inquiry*. 26: 571-613.

Jackobson, R. 1941. Kindersprache, Aphasie und Allemeine Lautgesetze, *Selected Writings* Vol. 1, 328-401. The Hague: Mouton, 1971.

Jespersen, O. 1904. *Lehrbuch der Phonetik*. Leipzig and Berlin.

Jones, C. 1978. Rounding and Fronting in Old English Phonology. *Lingua*. 46: 157-168.

Kaisse, E. 1992. Can [consonantal] Spread? *Language*. 68: 313-332.

Kawasaki, T. 1995. SV, [voice] and Coda Constraints. In L. Gabriele, D. Hardison, and R. Westmoreland (Eds.), *FLSM IV*. 27-38. Bloomington: The Indiana University Linguistics Club.

Kawasaki, T. 1996. Voicing and Coda Constraints. *McGill Working Papers in Linguistics*. 11: 24-45.

Kaye, J. D. 1990. 'Coda' Licensing. *Phonology*. 7: 301-330.

Kaye, J. D., J. Lowenstamm and J-R. Vergnaud. 1985. The Internal Structure of Phonological Elements: A Theory of Charm and Government, *Phonology Yearbook*. 2:305-328.

Kaye, J. D., J. Lowenstamm and J-R. Vergnaud. 1990. Constituent Structure and Government in Phonology. *Phonology*. 7: 193-231.

Kean, M. L. 1975. *The Theory of Markedness in Generative Grammar*. Ph. D. Dissertation, MIT.

Kenstowicz, M. 1994. *Phonology in Generative Grammar*. Cambridge: Blackwell.

Kiparsky, P. 1982. Lexical Phonology and Morphology. In I.-S. Yang (Ed.), *Linguistics in the Morning Calm*. Hanshin, Seoul: Linguistic Society of Korea.

Kiparsky, P. 1985. Some Consequences of Lexical Phonology. *Phonology Yearbook*. 2: 85-138.

Kisseberth, C. 1970. On the Functional Unity of Phonological Rules. *Linguistic Inquiry*. 1: 291-306.

Ladefoged, P. 1982. *A Course in Phonetics*. New York: Harcourt Brace Jobanovich.

Ladefoged, P. 1989. *Representing Phonetic Structure: UCLA Working Papers in Phonetics* 73. Los Angeles, California.

Leben, William. 1973. *Suprasegmental Phonology*. Ph. D. Dissertation, MIT.

Lamontagne, G. A. 1993. *Syllabification and Consonant Cooccurrence Conditions*. Ph. D. Dissertation, University of Massachusetts, Amherst.

Lombardi, L. 1991. *Laryngeal Features and Laryngeal Neutralization*. Ph. D. Dissertation, University of Massachusetts, Amherst.

Lombardi, L. 1995a. Laryngeal Neutralization and Alignment. In J. Beckman, L. W. Dickey, and S. Urbanczyk (Eds.), *Papers in Optimality Theory (University of Massachusetts Occasional Papers 18)*. 225-247.

Lombardi, L. 1995b. Laryngeal Neutralization and Syllable Wellformedness. *Natural Language & Linguistic Theory*. 13: 39-74.

Mascaró, J. 1983 . Phonological Levels and Assimilatory Processes. Paper presented at the 1983 GLOW Colloquium. York, England.

Mascaró, J. 1986. Compensatory Lengthening in Majorcan Catalan. In E. Sezer and L. Wetzels (Eds.), *Compensatory Lengthening* . Dordrecht: Foris.

Mascaró, J. 1987. A Reduction and Spreading Theory of Voicing and Other Sound Effects. Ms., Universitat Autònoma de Barcelona. [Published 1996 in *Catalan Working Papers in Linguistics* 4.2: 267-328].

McCarthy, J. 1979. *Formal Problems in Semitic Morphology and Phonology*. Ph. D. Dissertation, MIT.

McCarthy, J. 1988. Feature Geometry and Dependency: A Review. *Phonetica*. 43: 84-108.

McCarthy, J. and A. Prince. 1986. Prosodic Morphology. Ms., University of Massachusetts and Brandeis University.

McCarthy, J. and A. Prince. 1990. Prosodic Morphology and Templatic Morphology. *Perspectives on Arabic Linguistics II: Papers from the Second Annual Symposium on Arabic Linguistics. Current Issues in Linguistic Theory* 72. 1-54. Amsterdam and Philadelphia: John Benjamins.

McCarthy, J. and A. Prince. 1993. Generalized Alignment. *Yearbook of Morphology*. 79-154.

McCarthy, J. and A. Prince. 1994. The Emergence of the Unmarked: Optimality in Prosodic Morphology. In M. González (Ed.), *Proceedings of NELS 24*. 333-379. Amherst, MA: GLSA.

McCarthy, J. and A. Prince. 1996. Faithfulness and Reduplicative Identity. In J. Beckman, L. W. Dickey, and S. Urbanczyk (Eds.), *Papers in Optimality Theory (University of Massachusetts Occasional Papers 18)*. 249-384. Amherst: GLSA, University of Massachusetts.

Mills, E. 1984. *Senufo Phonology, Discourse to Syllable (A Prosodic Approach)*. Austin: The Summer Institute of Linguistics.

- Mills, R. 1975. *Proto South Sulawesi and Proto Austronesian Phonology*. Ph. D. Dissertation, University of Michigan.
- Mithun, M. and H. Basri. 1986. The Phonology of Selayarese. *Oceanic Linguistics*. 25: 210-54.
- Mohanan, K. P. 1983. The Structure of the Melody. Ms., MIT.
- Mohanan, K. P. 1993. Fields of Attraction in Phonology. In J. Goldsmith (Ed.), *The Last Phonological Rule: Reflections on Constraints and Derivations*. 61-116. Chicago: The University of Chicago Press.
- Mugane, J. and C. Gerfen. 1993. Aperture Geometry, Feature Cooccurrence and Meinhof's Law in Gikuyu. Paper presented at the Annual Conference in African Linguistics 24. Ohio State University.
- Murray, R. and T. Vennemann. 1983. Sound Change and Syllable Structure in Germanic Phonology. *Language*. 59: 514-528.
- Myers, S. 1987. Vowel Shortening in English. *Natural Language and Linguistic Theory*. 5: 485-518.
- Newton, B. 1972. *Cypriot Greek: Its Phonology and Inflections*. The Hague: Mouton.
- Osborn, H. 1948. Amahuaca Phonemes. *International Journal of American Linguistics*. 14: 188-190.
- Padgett, J. 1995. Feature Classes. In J. Beckman, L. W. Dickey, and S. Urbanczyk (Eds.), *Papers in Optimality Theory (University of Massachusetts Occasional Papers 18)*. 385-420. GLSA: University of Massachusetts, Amherst.
- Paradis, C. 1988a. On Constraints and Repair Strategies. *The Linguistic Review*. 6: 71-97.
- Paradis, C. 1988b. Towards a Theory of Constraint Violations. *McGill Working Papers in Linguistics*. 5: 1-43.
- Paradis, C. and J-F. Prunet. 1991. *Phonetics and Phonology 2: The Special Status of Coronals; Internal and External Evidence*. San Diego: Academic Press.
- Pater, J. V. 1996. *Consequences of Constraint Ranking*. Ph. D. Dissertation, McGill University, Montréal.
- Payne, D. L. 1981. *The Phonology and Morphology of Axininca Campa*. Dallas: Summer Institute of Linguistics.
- Piggott, G. L. 1991. Apocope and the Licensing of Empty-Headed Syllables. *The Linguistic Review*. 8: 287-318.

Piggott, G. L. 1992. Variability in Feature Dependency: The Case of Nasality. *Natural Language and Linguistic Theory*. 10: 33-77.

Piggott, G. L. 1994. Meinhof's Law and the Representation of Nasality, *MOT Phonology Workshop*. Toronto: Toronto Working Papers in Linguistics. 13 (1):123-145.

Piggott, G. L. and R. Singh. 1985. The Phonology of Epenthetic Segments. *Canadian Journal of Linguistics*. 30: 415-453.

Postal, P. 1968. *Aspects of Phonological Theory*. New York: Harper and Row.

Prince, A. and P. Smolensky. 1993. Optimality Theory: Constraint Interaction in Generative Grammar. Ms., Rutgers University and University of Colorado.

Pulleyblank, D. 1994. Neutral Vowels in Optimality Theory: A Comparison of Yoruba and Wolof. Ms., University of British Columbia.

Rehg, K. L. and D. G. Sohl. 1981. *Ponapean Reference Grammar*. Honolulu: The University Press of Hawaii.

Rice, K. 1992. On Deriving Sonority: A Structural Account of Sonority Relationships. *Phonology*. 9: 61-99.

Rice, K. 1993. A Reexamination of the Feature [sonorant]: The Status of 'Sonorant-Obstruents'. *Language*. 69: 308-344.

Rice, K. 1996. Default Variability: the Coronal-Velar Relationship. *Natural Language and Linguistic Theory*. 14:493-543.

Rice, K. and P. Avery. 1989. On the Interaction between Sonorancy and Voicing, *Toronto Working Papers in Linguistics* 10:65-82.

Rice, K. and P. Avery. 1991. On the Relationship between Laterality and Coronality. In C. Paradis and J-F. Prunet (Eds.), *Phonetics and Phonology 2: The Special Status of Coronals; Internal and External Evidence*. 101-124. San Diego: Academic Press.

Rivas, A. 1974. Nasalization in Guarani. *Papers from the 5th North East Linguistic Society*. 134-143.

Rose, S. 1996. Variable Laryngeals and Vowel Lowering. *Phonology*. 13: 73-117.

Rosenthal, S. 1989. *The Phonology of Nasal-Obstruent Sequences*. Masters thesis, McGill University, Montreal.

Rubach, J. 1996. Nonsyllabic Analysis of Voice Assimilation in Polish. *Linguistic Inquiry*. 27: 69-110.

Sagey, E. 1986. *The Representation of Features and Relations in Nonlinear Phonology*. Ph. D. Dissertation, MIT.

- Sakai, H. 1994. Alignment with Place Nodes: An Analysis of Lexical Domain Distinctions in Japanese. Paper presented at the WCCFL 13.
- Sapir, D. J. 1969. *A Grammar of Diola-Fogny: A Language Spoken in the Basse-Casamance Region of Senegal*. London: Cambridge University Press.
- Saussure, F. 1916. *Cours de linguistique générale*. (W. Baskin, Trans.). Paris: Payot.
- Scatton, E. A. 1984. *A Reference Grammar of Modern Bulgarian*. Columbus, Ohio: Slavica Publishers, Inc.
- Schane, S. 1984a. The Fundamentals of Particle Phonology. *Phonology Yearbook*. 1: 129-155.
- Schane, S. 1984b. Two English Vowel Movements: A Particle Analysis. In M. Aronoff and R. T. Oehrle (Eds.), *Language Sound Structure*. Cambridge: MIT Press.
- Schein, B. and D. Steriade. 1986. On Geminates. *Linguistic Inquiry*. 17: 691-744.
- Scobbie, J. 1991. *Attribute Value Phonology*. Ph. D. Dissertation, University of Edinburgh.
- Scobbie, J. 1992. Toward Declarative Phonology. In S. Bird (Ed.), *Edinburgh Working Papers in Cognitive Science*. 1-27. Edinburgh: University of Edinburgh.
- Selkirk, E. 1982. The Syllable. In H. van der Hulst and N. Smith (Eds.), *The Structure of Phonological Representations*. Vol. II, 337-383. Dordrecht: Foris.
- Selkirk, E. 1984. On the Major Class Features and Syllable Theory. In M. Aronoff and R. T. Oehrle (Eds.), *Language Sound Structure*. 105-136. Cambridge: The MIT Press.
- Shipley, W. F. 1956. The Phonemes of Northeastern Maidu. *International Journal of American Linguistics*. 22: 233-237.
- Shipley, W. F. 1963. *Maidu Texts and Dictionary*. Berkeley and Los Angeles: University of California Press.
- Sievers, E. 1885. *Grundzüge der Phonetik*. Leipzig: Breitkopf und Härtel.
- Singh, R. 1987. Well-Formedness Conditions and Phonological Theory. In W. Dressler et al. (Ed.), *Phonologica 1984*. Cambridge: Cambridge University Press.
- Smith, C and R Smith. 1971. Southern Barasano Phonemics. *Linguistics*. 75: 80-85.
- Stanley, R. 1967. Redundancy Rules in Phonology. *Language*. 43: 393-436.
- Steinbergs, A. 1985. The Role of MSC's in OshiKwanyama Loan Phonology. *Studies in African Linguistics*. 16: 89-101.

- Steriade, D. 1982. *Greek Prosodies and the Nature of Syllabification*. Ph. D. Dissertation, MIT.
- Steriade, D. 1995. Underspecification and Markedness. In J. Goldsmith (Ed.), *A Handbook of Phonological Theory*. 114-174. Cambridge: Blackwell.
- Tourville, J. 1991. *Licensing and the Representation of Floating Nasals*. Ph. D. Dissertation, McGill University, Montréal.
- Trigo, R. L. 1988. *On the Phonological Behavior and Derivation of Nasal Glides*. Ph. D. Dissertation, MIT.
- Vago, R. M. 1980. *The Sound Pattern of Hungarian*. Washington D. C.: Georgetown University Press.
- Vail, L. 1972. The Noun Classes of Ndali. *Journal of African Languages*. 11: 21-47.
- van der Hulst, H. 1989. Atoms of Segmental Structure: Components, Gestures and Dependency. *Phonology*. 6: 253-284.
- Watkins, L. J. 1984. *A Grammar of Kiowa*. University of Nebraska Press.
- Welmers, W. E. 1973. *African Language Structures*. Berkeley: University of California Press.
- Welmers, W. E. 1976. *A Grammar of Vai*. Berkeley and Los Angeles: University of California Press.
- Whitney, W. D. 1874. *Oriental and Linguistic Studies*. New York: Charles Scribner's Sons.
- Whitney, W. D. 1889. *Sanskrit Grammar*. Cambridge, MA: Harvard University Press.
- Williamson, K. 1987. Nasality in Ijo. In D. Odden (Ed.), *Current Approaches to African Linguistics*. Vol. 4, 397-415. Dordrecht: Foris.
- Wiswall, W. J. 1991. *Partial Vowel Harmonies as Evidence for a Height Node*. Ph. D. Dissertation, University of Arizona.
- Wonderly, W. 1951. Zoque II: Phonemes and Morphophonemes. *International Journal of American Linguistics*. 17: 105-123.
- Zec, D. 1988. *Sonority Constraints on Prosodic Structure*. Ph. D. Dissertation, University of Massachusetts, Amherst.
- Zoll, C. 1996. *Parsing Below the Segment in a Constraint Based Framework*. Ph. D. Dissertation, University of California, Berkeley.