

The Aeroacoustics of Nasalized Fricatives

by

Ryan Keith Shosted

B.A. (Brigham Young University) 2000
M.A. (University of California, Berkeley) 2003

A dissertation submitted in partial satisfaction of the
requirements for the degree of
Doctor of Philosophy

in

Linguistics

in the

GRADUATE DIVISION
of the
UNIVERSITY OF CALIFORNIA, BERKELEY

Committee in charge:
John J. Ohala, Chair
Keith Johnson
Milton M. Azevedo

Fall 2006

The dissertation of Ryan Keith Shosted is approved:

Chair

Date

Date

Date

University of California, Berkeley

Fall 2006

The Aeroacoustics of Nasalized Fricatives

Copyright 2006

by

Ryan Keith Shosted

Abstract

The Aeroacoustics of Nasalized Fricatives

by

Ryan Keith Shosted

Doctor of Philosophy in Linguistics

University of California, Berkeley

John J. Ohala, Chair

Understanding the relationship of aerodynamic laws to the unique geometry of the human vocal tract allows us to make phonological and typological predictions about speech sounds typified by particular aerodynamic regimes. For example, some have argued that the realization of nasalized fricatives is improbable because fricatives and nasals have antagonistic aerodynamic specifications. Fricatives require high pressure behind the supralaryngeal constriction as a precondition for high particle velocity. Nasalization, on the other hand, vents back pressure by allowing air to escape through the velopharyngeal orifice. This implies that an open velopharyngeal port will reduce oral particle velocity, thereby potentially extinguishing frication. By using a mechanical model of the vocal tract and spoken fricatives that have undergone coarticulatory nasalization, it is shown that nasalization must alter the spectral characteristics of fricatives, e.g. by reducing high-frequency energy and increasing the bandwidth of spectral prominences. These spectral modifications are liable to change the percept of fricatives at different places of articulation. It is hypothesized that nasalization generally has a deleterious effect on the acoustic distinctiveness of fricatives, explaining the typological rarity of nasalized fricatives. It also suggests that sibilant fricatives might be better at blocking the effects of nasal harmony.

John J. Ohala
Dissertation Committee Chair

Contents

List of Figures	iv
List of Tables	vi
1 Introduction	1
1.1 The Problem	1
1.2 Aeroacoustics of fricatives	4
1.3 Aeroacoustics of nasals	11
1.4 The Ohalian hypothesis considered	18
1.4.1 Ohala (1975)	19
1.4.2 Ohala and Ohala (1993)	20
1.4.3 Ohala, Solé, and Ying (1998)	21
1.4.4 Yu (1999)	22
1.4.5 Solé (1999)	23
1.5 Against the Ohalian hypothesis	24
1.5.1 Schadeberg (1982)	24
1.5.2 Gerfen (1999, 2001)	27
1.5.3 Coatzospan overview	28
Gerfen’s instrumental approach (1999, 2001)	29
Recommendations	30
1.6 Strong and weak versions of the hypothesis	34
1.7 Reports of nasalized fricatives	35
1.7.1 Applecross Scots Gaelic (Celtic, Scotland)	35
1.7.2 Chichimeco-Jonaz (Otopamean, Mexico)	37
1.7.3 Coatzospan Mixtec (Mixtecan, Mexico)	37
1.7.4 Epena Pedee (Choko, Colombia)	38
1.7.5 Ìgbò (Niger-Congo, Nigeria)	38
1.7.6 Icelandic	39
1.7.7 Inor (Semitic, Ethiopia)	40
1.7.8 Japanese	41
1.7.9 Umbundu (Niger-Congo, Angola)	41
1.7.10 Waffa (Papuan, Papua New Guinea)	41
1.7.11 Other ‘nasal harmonic’ languages	42

	TRANSPARENT fricative languages	43
	OPAQUE fricative languages	45
1.8	Summary	45
2	Method	48
2.1	Research hypotheses	48
2.2	Methodological overview	49
2.2.1	Spoken fricatives	49
2.2.2	Mechanical fricatives	50
2.3	Languages	50
2.4	Speakers	51
2.5	Stimuli	51
2.6	Spoken data	53
2.6.1	Audio	54
2.6.2	Oral flow	55
2.6.3	Nasal flow	56
2.6.4	Flow calibration	57
	Procedure	57
	Correlation coefficient	58
2.7	Mechanical fricatives	59
2.7.1	Model design	59
2.7.2	Model data	61
2.8	Acoustic analysis	62
2.8.1	Segmentation	62
2.8.2	Normalization	62
2.8.3	Zero-crossing rate	63
2.8.4	Power spectra	63
	Spectral averaging techniques	64
2.8.5	Parameterization of fricative spectra	66
	High frequency spectral slope (HiSlope)	68
	Low frequency spectral slope (LoSlope)	68
	Slope reference	69
	Dynamic amplitude (DynAmp)	70
	High wide-band frequency energy (HiBand)	70
	Spectral peak bandwidth	70
2.9	Flow analysis (spoken fricatives)	70
2.9.1	Segmentation	70
2.9.2	Normalization	71
2.9.3	Polynomial fitting	71
	Coefficients	71
	Correlation	71
	Norm of residuals	72
	Statistical evaluation of polynomial fit	72
2.9.4	Numerical integration	73
2.9.5	Maximal flow rate and flow rate at temporal center	75

2.10	Pressure analysis (mechanical fricatives)	75
2.11	Statistical Methods	75
2.11.1	Review of variables	75
	Continuous variables	75
	Categorical variables	75
2.11.2	Null hypotheses	76
	Spoken fricatives	76
	Mechanical fricatives	76
2.11.3	Linear statistical models	77
	Normality	77
	One-way analysis of variance	77
2.11.4	Non-linear models: Kruskal-Wallis	78
3	Results	79
3.1	Overview of the results	79
3.2	Spoken fricatives	79
3.2.1	Aerodynamic results	79
3.2.2	Acoustic results	84
3.3	Mechanical fricatives	90
3.3.1	Aerodynamic results	90
3.3.2	Acoustic results	92
	Ensemble-averaged data for mechanical fricatives	97
4	Discussion and Conclusions	100
4.1	Summary of the results	100
4.2	Nasal harmony	104
4.3	Velopharyngeal dysfunction	105
4.4	Voiceless nasals	106
4.5	Sibilants and non-sibilants	108
4.6	Universals, rarities, and the expanding IPA	109
4.6.1	An infinite phonetic alphabet?	110
4.6.2	The IPA as a Cartesian coordinate system	113
4.6.3	Nasalized fricatives: Shaded or empty cell?	114

List of Figures

1.1	Tube model of the vocal tract	1
1.2	Relationship of pressure behind a constriction in a tube and volume velocity	7
1.3	FFT of a uniformly-distributed random process	9
1.4	FFT of a voiceless alveolar fricative	10
2.1	Audio, oral flow, and nasal flow during $\tilde{a}f\tilde{a}$ (Hindi)	54
2.2	FFT of an alveolar [s] produced with speaker wearing Scicon OM-2 (oral mask).	56
2.3	FFT of an alveolar [s] produced without the Scicon OM-2 oral mask.	57
2.4	Area function of an American English alveolar fricative	60
2.5	Pressure and audio output of mechanical fricative	62
2.6	Nine windows arrayed across a fricative signal	64
2.7	Center frame acoustic data	65
2.8	Application of Hamming function	66
2.9	DFT of velar fricative between two nasal vowels	67
2.10	Spectral parameterization	68
2.11	Slope diagram	69
2.12	Nasal flow during labiodental fricative between nasal vowels	72
2.13	Oral flow during labiodental fricative between nasal vowels	73
2.14	Residuals for a cubic polynomial fitted to nasal flow	74
3.1	Boxplots of integrated nasal flow by nasal context	81
3.2	Boxplots of nasal flow maxima by nasal context	82
3.3	Boxplots of nasal flow at temporal center of fricative by nasal context	82
3.4	Boxplots of integrated oral flow by nasal context	83
3.5	Boxplots of oral flow maxima by nasal context	83
3.6	Boxplots of oral flow at temporal center of fricative by nasal context	84
3.7	Boxplots of oral flow maxima by V1	85
3.8	Boxplots of $\bar{F}$ by fricative (place of articulation)	86
3.9	Boxplots of Zero-crossing rate by fricative (place of articulation)	87
3.10	Boxplots of HiSlope in the first fricative frame by nasality condition of V1	88
3.11	Boxplots of HiBand measures in the first fricative frame by nasality condition of V1	89
3.12	Boxplots of spectral peak bandwidth (Hz) in by nasality condition of V1	90

3.13	Aerodynamic performance of the fricative model	91
3.14	Averaged spectra of mechanical alveolar fricatives with VPO = 0 and 0.713 cm ²	93
3.15	$\bar{F}$ for mechanical alveolar fricatives with differing VPOs	94
3.16	HiSlope for mechanical alveolar fricative with differing VPOs	95
3.17	HiBand for mechanical alveolar fricative with differing VPOs	96
3.18	LoSlope for mechanical alveolar fricative with differing VPOs	97
3.19	DynAmp for mechanical alveolar fricative with differing VPOs	98
3.20	Spectral peak bandwidth for mechanical alveolar fricative with differing VPOs	99
4.1	Photograph of the mechanical fricative model	116

List of Tables

1.1	Fricatives through which nasalization ‘spreads’ in Coatzospan Mixtec	28
1.2	Fricatives that block nasalization in Coatzospan Mixtec	29
1.3	Nasalized fricatives of Applecross Scots Gaelic	37
1.4	Nasal and oral fricatives in Ìgbò	38
1.5	<i>Constrictives nasales</i> in Icelandic	40
1.6	Nasal contrasts in Waffa	42
1.7	Nasal harmony definitions	43
1.8	Type IV nasal harmony languages	44
1.9	Type V nasal harmony languages	47
2.1	Continuous acoustic variables	75
2.2	Continuous aerodynamic variables	75
2.3	Categorical variables	76
3.1	Numbers of fricatives analyzed aerodynamically	80
3.2	Mean values for aerodynamic measures	80
3.3	ANOVA results for nasal aerodynamic measures by nasal context	81
3.4	ANOVA results for oral aerodynamic measures by nasal context	81
3.5	Tukey’s HSD for pressure	92
3.6	$\bar{F}$ of mechanical fricatives at differing velopharyngeal openings	94
3.7	Zero-crossing rate of mechanical fricatives at differing velopharyngeal openings	94
3.8	HiSlope of mechanical fricatives at differing velopharyngeal openings	95
3.9	HiBand of mechanical fricatives at differing velopharyngeal openings	95
3.10	LoSlope of mechanical fricatives at differing velopharyngeal openings	96
3.11	DynAmp of mechanical fricatives at differing velopharyngeal openings	96
3.12	Spectral peak bandwidth of mechanical fricatives at differing velopharyngeal openings	97
3.13	Frame-by-frame variation in mechanical fricatives	98

Acknowledgments

I wish to thank Professors John Ohala, Keith Johnson, Milton Azevedo, and Ian Maddieson for their critical sense and generous guidance throughout the course of this project. I am also grateful to Ronald Sprouse, who designed the Matlab routines for data collection. Many individuals have given feedback when I have presented parts of this material at conferences; to them I also express gratitude. Despite the contributions of others, I alone am responsible for any errors, whether of fact or interpretation, that may persist in the manuscript.

Chapter 1

Introduction

1.1 The Problem

When the human vocal mechanism is reduced to a simple model of conjoined tubes (see Figure 1.1), certain mechanical properties of the system can be derived. It becomes clear that the properties of the system constrain its output (the sounds the system can emit). While there are many constraints on the vocal mechanism of humans, this study will focus on a single aerodynamic constraint that has a growing importance in the phonetic and phonological literature, viz. that nasalization and oral¹ frication cannot be produced simultaneously.

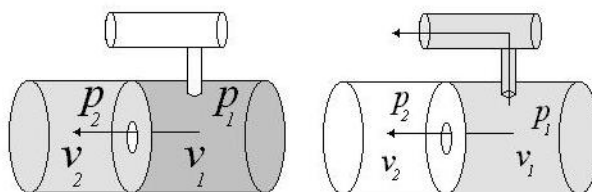


Figure 1.1: Tube model of the vocal tract.

¹By ‘oral’, I refer to a place of articulation anterior to the velopharynx, more specifically, ‘buccal’. There is no reason to doubt that glottal or pharyngeal fricatives may be nasalized, i.e. $[\tilde{h} \tilde{ɦ} \tilde{ħ} \tilde{ç}]$.

From a mechanical standpoint, it seems clear that nasal and oral fricative sounds have antagonistic aerodynamic specifications. These seem to preclude them from being produced at the same time in the same vocal tract. Oral fricatives require high pressure behind a constriction in order to achieve high particle velocity, itself a determiner of the aperiodic noise characteristic of fricative acoustics. At the same time, nasals require an open velopharyngeal orifice, which vents back pressure. While fricatives are the present object of study, it is worth noting that no one has ever claimed that a language has stops produced with a lowered velum.² The burst characteristics of simple stops and affricates (burst plus frication) are predicated on pressure build-up. If the antagonism of simultaneous pressure build-up and pressure leakage obviates nasal stops and affricates, then it should also obviate nasalized fricatives. However, this kind of apagogical argument does little to answer the numerous reports of nasalized fricatives in the world's languages (see Sections 1.7.9, 1.7.3, and 1.7).

Logically, under the assumption of constant transglottal flow, pressure behind the constriction and particle velocity across the constriction must be sacrificed during a nasalized fricative. The questions of whether and to what extent this sacrifice may be 'fatal' to the fricative will occupy the present thesis. In terms of phonological systems and typological patterns, the aeroacoustic sacrifice represented by the nasalization of fricatives may have a number of outcomes, all of which are empirical issues to be addressed presently:³

1. Nasalized fricatives are not found in the languages of the world;
2. Fricatives prevent the 'spread' of nasalization in nasal harmony systems;
3. Some fricatives are more likely to nasalize than others based on their aeroacoustic properties.

One of the goals of laboratory phonology is to make sensible predictions about sound systems based on empirical evidence. To prove conclusively that the laws of aerodynamics and the unique geometry of the human vocal tract rule out the existence of nasalized fricatives and, moreover, that nasalized fricatives are unattested in the languages of the

²Ladefoged and Maddieson (1996: 134) observe that such sounds can only be produced behind the opening to the nasal cavity, e.g. Sundanese [ʔ].

³An additional outcome, not addressed in any detail here, is the emergence of transitional segments, sometimes epiphenomenal stops, at the juncture of a nasal consonant and an oral fricative, e.g. team[p]ster. For further discussion, see Ali et al. (1979), Fourakis and Port (1986), Ohala (1995).

world would be a boon to the laboratory approach. However, as is often the case in science, the waters are a bit murkier than that. Various authors have challenged the universalist laboratory-influenced claims by positing nasalized fricatives in a variety of geographically- and typologically-diverse languages, though often with unsatisfactory documentation. It is beyond the scope of the present study to evaluate the empirical basis of these claims, though they will be reviewed in detail (see Sections 1.5.1, 1.5.2, and 1.7). Instead, the present study will address the acoustics of nasalized fricatives, an issue that seems like a logical extension of the controversy. Rather than asking whether or not nasalized fricatives exist in the languages of the world (an empirical task, which to be undertaken in any comprehensive fashion, would involve collecting aerodynamic evidence on at least four continents) the present study investigates the spectral characteristics of nasalized fricatives. If such sounds are possible, what might they sound like?

To answer the question, one might consider three different kinds of sounds:

1. Nasalized fricatives of a language which is reported to have them or in which phonological nasal harmony is likely to give rise to them. Such languages are of three classes:
 - (a) Languages like Waffa (Papuan, Papua New Guinea)⁴ in which nasal fricatives are simply posited as part of the phonological inventory, without reference to nasal harmony (Stringer and Hotz 1973).
 - (b) Languages like Applecross Scots Gaelic, in which nasal harmony operates ‘through’ fricatives and explicit claims have been made regarding the fricatives’ nasalized status (Ternes 1989).
 - (c) Languages like Apinayé (Ge, Brazil) in which ‘nasal harmony’ or ‘nasal spreading’ is reported to operate ‘through’ fricatives, so fricatives between nasal segments may *potentially* be nasalized. (Walker 2000: 66)⁵
2. Fricatives produced (by speakers of any language with fricatives and phonemic nasals)

⁴Throughout, I will include information about the family and primary national affiliation of understudied languages. Hence, the parenthetical comment (Papuan, Papua New Guinea) indicates that Waffa is a Papuan language spoken in Papua New Guinea; Apinayé is a Ge language spoken principally in Brazil, and so forth. In cases where this extra information is extant in the language name, e.g. ‘Applecross Scots Gaelic’, the genetic and national information is not provided. All genetic classifications and national affiliations come from Gordon (2005).

⁵For such languages, I must emphasize that there is *no explicit claim* that fricatives become nasalized. It is only a possibility. See Section 1.7.11 for further discussion.

in environments where they are likely to undergo some degree of coarticulatory nasalization, e.g. C in $\tilde{V}C\tilde{V}$ strings;

3. ‘Fricatives’ (literally hissing or hushing noises) produced by a mechanical model of the vocal tract in which pressure can be systematically vented to replicate the effects of nasalization.

In the present study, only the last two types of sounds will be collected and analyzed. It is ultimately concluded that nasalization indeed changes the spectral characteristics of fricatives in certain ways. However, at present the question of their perceptibility will remain the object of conjecture rather than rigorous investigation. It is hoped that the present study will contribute fundamental aerodynamic and acoustic data to the study of ‘nasalized fricatives’ and that it will also lead to discussions of ‘fricative space,’ i.e. the dimensions along which fricatives are perceptually categorized and managed in phonological inventories.

1.2 Aeroacoustics of fricatives

During inspiration, air flows into the respiratory system (through the mouth or nose) because the alveolar (lung) pressure is less than the pressure at the mouth and/or nose (i.e. atmospheric pressure). The decrease in pressure is motivated by the upward and outward movement of the ribs, along with the downward movement of the diaphragm, which enlarges the thoracic cavity and hence, lung volume (Cotes et al. 2006: 99). Conversely, during expiration, air flows out of the respiratory system because alveolar pressure exceeds atmospheric pressure. The physiological mechanism for increasing lung pressure is the relaxation of the inspiratory muscles and subsequent elastic recoil of the lung tissue.

According to Boyle’s law, “at a constant temperature a gas volume is inversely related to its pressure,” or $PV = k$ where P is pressure, V volume, and k a constant (Cotes et al. 2006: 57). When the gas volume—the amount of space a gas occupies—is decreased, the pressure increases, and vice versa.

The movement of air between two regions, e.g. the atmosphere and the lungs, is conditioned by the difference in pressure between the two. Specifically, air will flow from a region of relatively high pressure to one of relatively low pressure. As this difference in pressure Δp increases, the flow rate or volume of air per unit time U will also increase.

However, when air moves at sufficiently different velocities through an airway, different equations are necessary to express relationships between pressure and flow. This is due to resistance, or the friction that individual molecules encounter as they pass through the airway.

When air flows at high velocities, especially through a conduit with irregular walls, the flow is generally disorganized, even chaotic, and tends to form vortices and eddies that interact with each other in unpredictable ways. This is called turbulent flow. Because of the relatively greater resistance encountered by individual molecules in turbulent flow, it requires more energy for a specific quantity of molecules to pass in a given unit time. In fact, to double the volume of gas per unit time (or flow rate) U one must quadruple the driving pressure Δp according to Equation 1.1, where Δp is the difference in pressure between two points and U_t is the volume velocity for turbulent flow (Daugherty and Franzini 1965). This equation presupposes that the radius of the airway is held constant.

$$\Delta p = kU_t^2 \quad (1.1)$$

At lower velocities, vortices tend not to form, so the individual molecules move in relatively straight lines and experience less resistance.⁶ When these conditions obtain, the flow is called laminar. Unlike turbulent flow, where Δp must be quadrupled in order to achieve a doubling of U , laminar flow rate U_l is directly proportional to the driving pressure. Accordingly, to double the flow rate U_l , one need only double the driving pressure Δp . Known as Poiseuille's Law, Equation 1.2 is said to govern laminar flow; η is the gas viscosity, ℓ is the length of the tube, and r is the radius (Cotes et al. 2006: 152).

$$\Delta p = U_l(8\eta\ell/\pi r^4) \quad (1.2)$$

In effect, this means that if the radius for laminar flow is doubled, all else being equal, the resistance decreases sixteen times. For turbulent flow, for any particular flow rate, the pressure drop is dependent on the fifth power of the radius of the tube (the Fanning equation) (Daugherty and Franzini 1965):

⁶In a laminar flow through a tube, the flow can be visualized as a series of concentric cylinders, each moving at a different velocity. The cylinder of air closest to the wall of the tube has the lowest velocity; this value gradually increases towards the center of the tube. Hence, if the leading particles in each concentric cylinder were viewed in profile, together they would appear as an advancing parabola with the fastest moving particle at the vertex.

$$\Delta p \propto \frac{1}{r^5} \quad (1.3)$$

Fricatives are produced under a turbulent airflow regime, so Equations 1.1 and 1.3 apply to sounds like [s f x] and, to a lesser extent, to sounds like [z v ʃ]. A fricative is said to occur in the vocal tract when a fast-moving jet strikes an obstacle (which need not be perpendicular to the flow) or moves through a channel that narrows and expands abruptly (Johnson 1997). The air that emerges from the constriction or passes the obstacle expands and forms a turbulent jet, producing noise (Shadle 1997: 44). To understand how the appropriate velocity is achieved, it will be necessary to review a number of aerodynamic principles and their significance for sound production.

Assuming no work, heat transfer, or change of elevation between two points in a tube, 1 and 2, a form of Bernoulli's equation can be derived to relate the pressure and velocity at those same two points.

$$-gH_L = \frac{p_1 - p_2}{\rho} + \frac{v_2^2 - v_1^2}{2} \quad (1.4)$$

This equation formalizes the relationship between p , particle velocity v , cross-sectional area A (at points 1 and 2), gravitational acceleration g , head loss H_L , and volume velocity U (Shadle 1997). Using the relation of volume velocity U to particle velocity v , $U = vA$, along with the assumption that U will be the same at any point along the duct and assuming $H_L = 0$ (i.e. the flow is frictionless before reaching point 2), we can rearrange the variables as in Equation 1.5.

$$U = \frac{cd \cdot A_1}{\sqrt{1 - (A_2/A_1)^2}} \sqrt{\frac{2(p_1 - p_2)}{\rho}} \quad (1.5)$$

Where ρ is the fluid density ($=1.139 \text{ kg/m}^3$); p_2 is atmospheric pressure ($=1.01325 \times 10^5 \text{ pa}$); cd is a dimensionless discharge coefficient; A_1 is the cross-sectional area of the orifice ($=0.1 \text{ cm}^2$); A_2 is the cross-sectional area of the duct ($=10 \text{ cm}^2$); and p_1 varies above atmospheric pressure p_2 . The value of cd depends on the Reynolds number (quite low in this case) and the ratio of the orifice to pipe diameter. Based on the discharge coefficient function found in Doebelin (1983) and cited by Shadle (1997), $cd = 0.6$ for present purposes. The measurement of a typical fricative constriction, A_1 comes from Shadle (1997: 44). The volume velocity output (in m^3/s) of Equation 1.5 is shown in Figure 1.2.

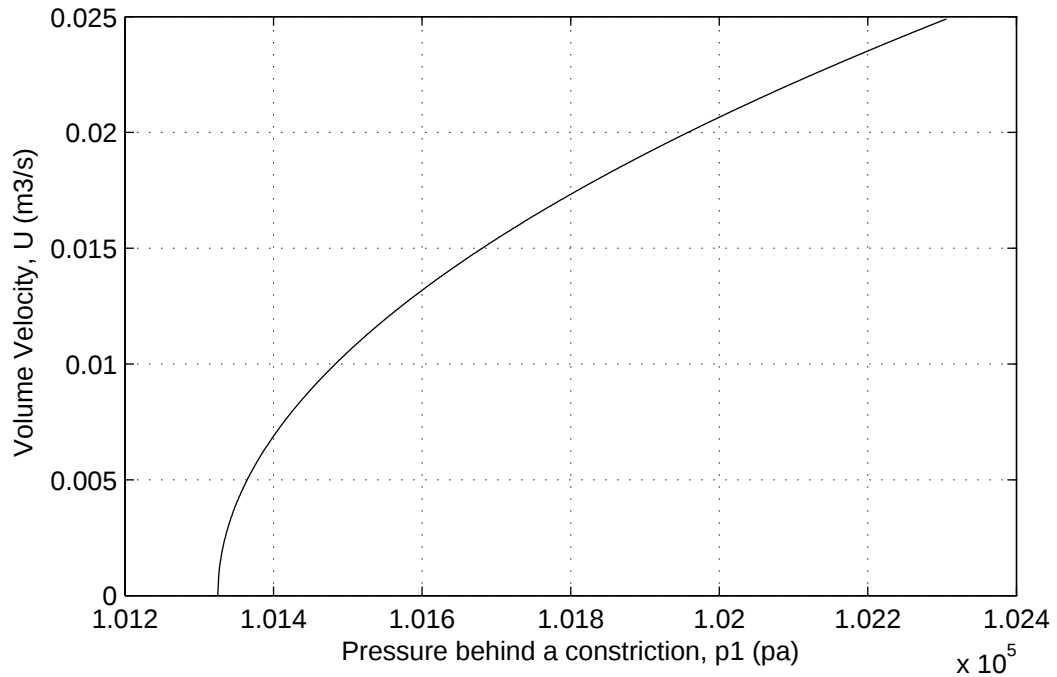


Figure 1.2: Relationship of pressure behind a constriction p_1 and volume velocity U as expressed in Equation 1.5. Atmospheric pressure, $p_2 = 1.01325 \times 10^5$ pa, so volume velocity becomes positive only after p_1 increases beyond this point.

The shape of the curve is logarithmic. An increased pressure gradient Δp or $p_2 - p_1$ produces higher volume velocity. Since p_2 or atmospheric pressure will not generally change during speech production, it is presumably safe for us to base the variation in volume velocity on p_1 , or the pressure behind the constriction. For present purposes, p_1 can be said to occur on the upstream side of an oral constriction and p_2 on the downstream side. This is what might be expected during the articulation of an oral fricative like [s], where the downstream pressure is low with regard to p_1 , the pressure behind the lingual constriction. The equations above indicate that as the pressure behind the constriction increases, the volume velocity U increases logarithmically.

All of this has important ramifications for the production of obstruents in general and fricatives in particular. As Δp increases, e.g. by the increase of p_1 (assuming constant p_2), v and U also increase. When the resultant high-velocity jet strikes an obstacle like the teeth or alveolar ridge, the turbulence of the airstream is magnified, creating more vortices.⁷

⁷If certain conditions obtain (based on jet thickness, jet standoff distance, and flow rate), a “sinuosity in

According to Gibson (1999: 83) turbulent flows are “dominated by a nonlinear force that randomly scrambles the motion on all length scales permitted by other forces that tend to damp out the turbulence.” The dynamics of turbulence are illustrated by the following equation where t is time, $\bar{\omega} = \text{curl}$, $\bar{v}$ is the vorticity, $\bar{\tau}$ is the viscous stress tensor, and the density ρ is assumed constant.

$$\frac{\partial \bar{v}}{\partial t} = \bar{v} \times \bar{\omega} - \nabla B + \nabla \cdot (\bar{\tau} / \rho) \quad (1.6)$$

Moreover, the Bernoulli Group $B = v^2/2 + p/\rho + gx_3$, where p is the pressure, g is the gravity, and x_3 is up. Turbulence occurs when the inertial-vortex forces ($\bar{v} \times \bar{\omega}$) per unit mass exceed the viscous forces $\nabla \cdot (\bar{\tau} / \rho)$. The ratio of inertial forces to viscous forces is the Reynolds number⁸ $Re = UL/v$, where U is a characteristic velocity and L is a characteristic length scale for the flow. Based on these equations, the definition of turbulence given by Gibson (1999) is this:

[A]n eddy-like state of fluid motion where the inertial-vortex forces are larger than any other forces that tend to damp them out.

In its first stages of development in a flow, turbulence appears as viscous eddies forming on the boundary layers of solid surfaces (e.g. the boundary layer around the teeth or alveolar ridge). These tend to break up into more random eddies as the jet of fast-moving fluid continues to interact with the slow-moving boundary layer and the eddies that emanate from it.

It is the randomness of turbulent flow that causes the ‘random’ high-frequency energy typical of fricatives. It is important to note, however, that the oft-cited acoustic ‘randomness’ of natural fricatives is far from random in any mathematical sense. This can be illustrated by simply computing the frequency content (or Fast Fourier Transform) of a computer-generated, uniformly-distributed random process (i.e. white noise), as shown in Figure 1.3. The unique resonant properties of a natural fricative are easily observed when

the growing wave” will develop, allowing for the possibility of a so-called ‘whistled fricative’ (Coltman 1968, Shadle 1983, Shosted 2006b).

⁸Reynolds numbers Re of less than 100 are associated with completely laminar flow; $Re > 10,000$ is associated with fully turbulent flow. During quiet breathing in the trachea, $Re = 1500$ so the flow is characterized as ‘partly turbulent’ (Cotes et al. 2006: 153). A comparable (or higher) intermediate value is likely during speech, which can thus be considered ‘partly turbulent’, as well. Only under conditions where the Reynolds number is extremely low can it be said that viscosity plays an important role in fluid dynamics, so the Reynolds number (and hence, viscosity) are probably of little relevance during the production of speech.

the spectrum of a fricative is compared with the spectrum of computer-generated white noise. The spectrum of a natural alveolar fricative, uttered by the author, is shown in Figure 1.4.

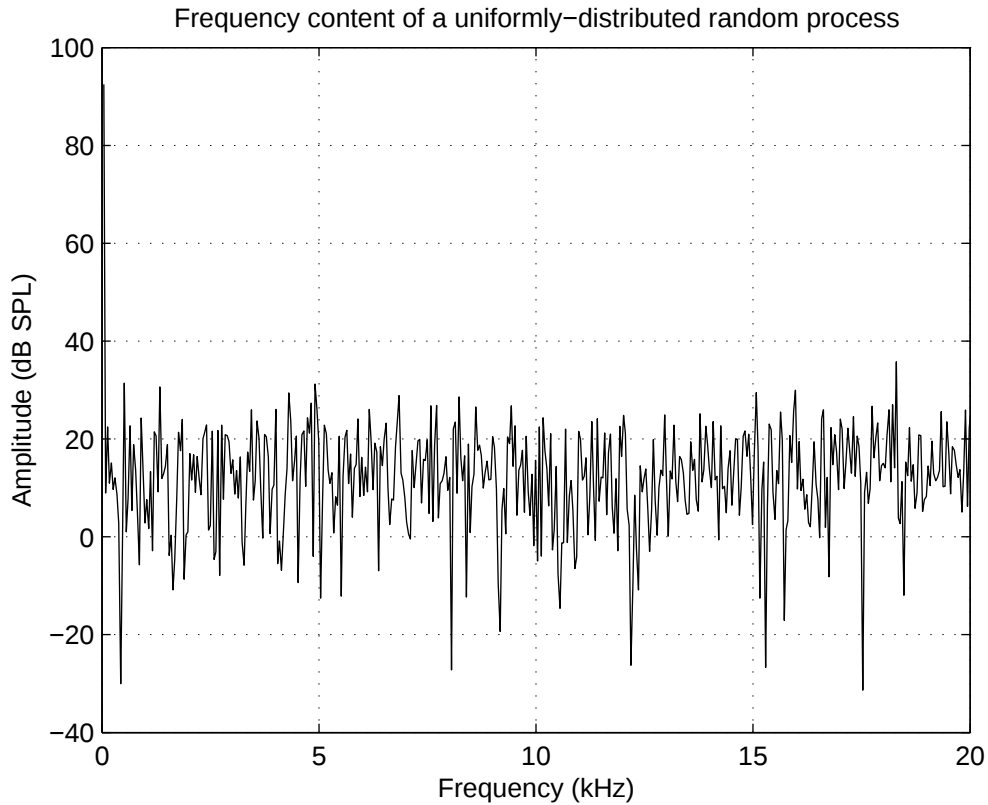


Figure 1.3: FFT of a computer-generated, uniformly-distributed random process. The lack of spectral peaks or formants demonstrates the truly random nature of the original signal. When compared with the spectrum of [s] presented in Figure 1.4, it becomes obvious that a natural fricative, with specific formants, is not truly a random signal.

For the computer-generated signal, there appears to be roughly equal power at any center frequency, having a given bandwidth. This is clearly not the case for the natural fricative. Here, there are peaks and valleys in spectral energy, indicating that the noise produced during [s] is not mathematically random.

The spectral prominences in the natural fricative are caused largely by the resonance cavity ‘downstream’ of the oral constriction (in this case, the lingual constriction formed at the alveolar ridge).⁹ The spectral envelope—or spatial configuration of these

⁹The size of the constriction determines whether or not (and to what extent) the upstream resonator

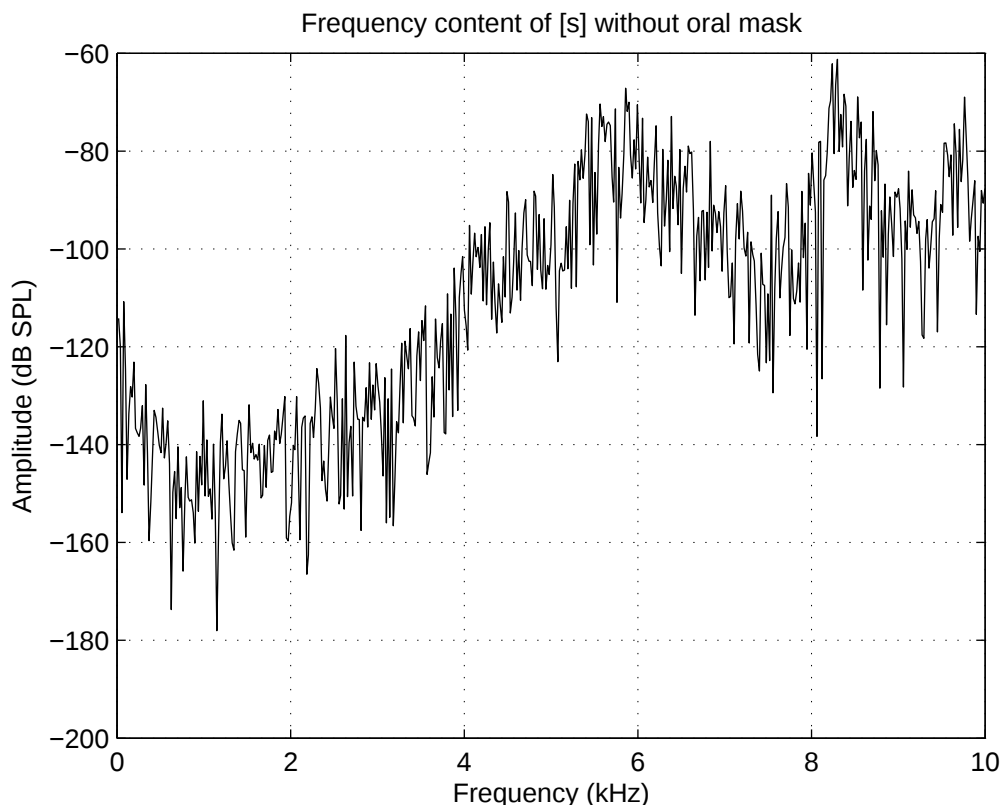


Figure 1.4: FFT of a naturally-produced [s], uttered by the author. The natural peaks in the signal contrast with the flat spectrum produced by a mathematically random process, e.g. Figure 1.3.

peaks—will vary for fricatives with different places of articulation because the dimensions of the resonating cavity vary for each. Because the resonating cavity is very small or non-existent for labiodental and bilabial fricatives, the spectra of sounds such as [f] and [ɸ] tend to be relatively flat (though presumably, not as flat as the white-noise spectrum in Figure 1.3).

Stevens (1998: 103) explains that sound is generated by turbulence at a surface (e.g. the palate for a velar fricative) or obstacle (e.g. the upper incisors for an alveolar fricative) in the vocal tract. He claims that the sound may be concentrated in a narrow region or may be distributed over a region that extends up to one centimeter downstream of the constriction.

may play a role (Stevens 1998: 141–142).

Shadle (1985) has provided experimental results demonstrating that sound power generated in the middle- and high-frequency range by this kind of turbulent flow (while constriction size is maintained constant) is proportional to the sixth power of the velocity of the air flow. Because we know the velocity of the flow is proportional to $\sqrt{\Delta P}$ where ΔP is the pressure drop across the constriction, we know that the sound power generated by a turbulent noise source is proportional to ΔP^3 (Stevens 1998). Though with some variation, this relationship between sound power and pressure drop has been observed experimentally by Hixon et al. (1967) and Badin (1989) among others. Moreover, since the radiated sound pressure is proportional to the square root of the sound power, we can figure the magnitude of the sound pressure source to be proportional $\Delta P^{3/2} A^{1/2}$ (where A is the cross-sectional area of the constriction). Finally, as Goldstein (1976) has shown, the spectrum of the sound pressure source resulting from obstacle-generated turbulence usually has a broad peak at a frequency proportional to u/d where u is the velocity of the airstream and d is the cross-dimension of the constriction.

This suggests that, all else being equal, lowering the pressure behind the constriction will decrease the radiated sound pressure generated by the turbulent noise (Stevens 1998, Shadle 1985) and lower the frequency at which the spectral peak will occur (Goldstein 1976). For example, if the pressure behind the alveolar constriction during the production of [s] were decreased by some factor ζ , then the sound pressure level generated during the production of this particular [s] would decrease by some multiple of ζ . This multiple, it is safe to say, will be determined by the cross-dimension of the constriction and the configuration of jet exit and obstacle (if there is one).

1.3 Aeroacoustics of nasals

In articulatory terms, nasalization may occur whenever the palatine aponeurosis, or soft palate, descends into the oropharynx (Bell-Berti 1993). With the soft palate lowered, if a standing wave is generated in the vocal tract, usually through the rapid vibration of the vocal folds, acoustic resonance is said to take place in the nasal passage.¹⁰ Acoustic nasalization, however, is only detectable at certain levels of velopharyngeal aperture. Hence, nasalization is a gradient phenomenon, both in articulatory and perceptual

¹⁰The tube produced by the opening of the velopharyngeal port is technically a resonator, regardless of the presence of a noise source.

terms (Beddor 1993). In a study of “hypernasal” speech among speakers with varying degrees of velopharyngeal inadequacy (i.e. cleft palate) Warren et al. (1993: 143) concluded that listeners “usually perceive hypernasal resonance when the velopharyngeal opening for nonnasal consonants is greater than 0.10 cm^2 [10 mm^2], and there is almost always some hypernasality perceived when the opening is greater than 0.2 cm^2 [20 mm^2].” He goes on to observe that while “the amount of opening into the nasal cavity influences the degree of perceived hypernasality, other factors such as status of the nasal airway and placement of oral structures also affect the perceptual outcome.”¹¹

Sounds emanating from the human vocal tract are acoustic structures based on the natural frequencies at which air vibrates in the tract. Because articulators are dynamic and can be repositioned in a variety of ways, the geometry of the tract and therefore the natural frequencies at which the air vibrates can change significantly. For example, a constriction in the pharynx (which falls near a pressure maximum in the standing wave of the first resonant frequency) increases the amplitude of that frequency for the low vowel [a] (vis-à-vis the same frequency for a constrictionless vocal tract configuration, i.e. [ə]) (Chiba and Kajiyama 1941). Frequencies with relatively prominent amplitudes, known as formants or “poles” in terms of complex analysis,¹² are the typical acoustic output of a tube with no sidebranches. When the velum is lowered, however, this classical “one-tube” model is fundamentally altered. For nasal consonants, with complete oral occlusion, the principle “tube” for the generation and emission of sound extends from the glottis to the nares (nostrils). There is, however, a significant sidebranch to this naso-pharyngeal passage, viz. the oral cavity. This is also true during the production of nasal vowels, with the complication that the oral cavity acts as an escape valve for the transglottal flow. As Fant (1970) has described, the additional side branch contributes antiformants or “zeros” in terms of complex number theory.¹³ Thus, the source of acoustic complexity in nasals (the appearance of oral antiformants in addition to naso-pharyngeal formants) is also their

¹¹Maeda (1993: 148–149) remarks that “[v]elar lowering not only opens the port, it also modifies the area function in the vicinity of the passage to the oral cavity. Experimenting with an analog model, House and Stevens (1956) concluded that the oral cavity area change contributed only a minor spectral modification. Using an articulatory synthesizer, Bell-Berti and Baer (1983) also demonstrated negligible effects of the area change, although they included the oral tract area modification in their simulation experiments. It is not unreasonable, therefore, to model the velopharyngeal port opening without changing oral cavity area function.”

¹²In the branch of mathematics investigating complex numbers (e.g. $a + bi$), a pole of a holomorphic function is a certain type of simple singularity that behaves like the singularity $\frac{1}{z^n}$ at $z = 0$. A pole of the function $f(z)$ is a point $z = a$ such that $f(z)$ approaches infinity as z approaches a .

¹³A zero of a holomorphic function f is a complex number a such that $f(a) = 0$.

definitive attribute (Kurowski and Blumstein 1993: 198).

In the simplest of terms, antiformants are components of a sound that cancel out other components. The situation is essentially one of direct and reflected waves, where the “direct waves” resonate in the naso-pharynx and the “reflected waves” resonate in the oral cavity Johnson (1997: 149). The reflected waves of the oral cavity have exactly the same phase as the direct waves of the naso-pharynx and therefore cancel out specific frequency components of the sound that is emitted from the nose. This is clearly the case for nasal consonants, where the oral cavity is sealed at one end and the standing wave patterns in the oral sidebranch interfere with and cancel out specific frequencies associated with the standing wave patterns of the naso-pharyngeal acoustic signal.¹⁴ In the case of nasalized vowels, nasal glides, and nasal fricatives, the oral occlusion may only approach 0 cm², so it seems reasonable to suggest that the antiformants in the spectra of these kinds of sounds are relatively less influential than the antiformants¹⁵ in the spectra of nasal “stops.”

This leads us to a closer consideration of the definition of “nasal consonant.” The canonical definition follows Stevens (1998: 305): “A nasal consonant is produced with a velopharyngeal opening but with a complete closure of the main vocal tract at some point within the oral cavity.” This is clearly the case for such consonants as [m ɱ n ɲ ŋ ɳ] but fails to describe classes of nasal consonant other than stops, e.g. glides like [ĩ ẽ ỹ] and nasal fricatives like [ʃ̃ ʒ̃ ʒ̃̃]. During the articulation of nasal glides and putative nasal fricatives, air is discharged from both the nose and the mouth, so it seems reasonable to group these sounds with the class of nasal vowels. Instead of modeling them with two conjoined tubes, one open and one closed (as is traditionally done for nasal stops), nasal vowels, nasal glides, and nasal fricatives should be modeled by two conjoined tubes that are both open to the atmosphere. Accordingly, the following discussion will concentrate on the acoustics of nasal vowels, not nasal stops, as the closest analog to the acoustics of nasal fricatives.¹⁶

¹⁴The situation is further complicated by the fact that the naso-pharyngeal and oro-pharyngeal tracts each contribute their own resonant frequencies; to the extent that the resonant frequencies are the same, they will cancel each other out.

¹⁵It is also worth noting that antiformants arising from the nasal sinuses play an important role in the acoustic analysis of nasals, as shown by Fujimura (1962). Similarly, the piriform sinus (also known as the *pyriform fossa*, the narrow tube above the vocal folds bounded by the epiglottis and aryepiglottic rim) contributes a zero in the speech spectrum—at around 4kHz (Dang et al. 1995, Sundberg 1972). Hence, both nasal and oral antiformants have spectral importance in speech.

¹⁶This, however, is not an unproblematic model. With regular oral fricatives, it is assumed that the constriction is usually close enough (at least for anterior fricatives) that significant acoustic coupling of the front and back cavities may be disregarded (Stevens 1998, Johnson 1997). For nasalized fricatives, however, the properties of the ‘back’ cavity must be taken into account since this includes the velopharyngeal port.

The determination of the locations of poles and zeros for nasal vowels is a rather complicated enterprise, as the geometries of the tract are difficult to pin down for individual speakers. Most studies (Delattre 1954, House and Stevens 1956, Hattori et al. 1958, Fant 1970, Fujimura and Lindqvist 1971, Bell-Berti and Baer 1983, Hawkins and Stevens 1985, Bognar and Fujisaki 1986) indicate that during nasalization there is a relative weakening of the first formant peak and a variety of secondary cues, such as a relative strengthening of the spectrum in the vicinity of 250 Hz.

In addition to the appearance of antiformants in the sound spectrum (Fujimura 1962) nasalization also tends to widen formant bandwidths (Johnson 1997, Stevens 1998). Stevens (1998: 310) observes that this is due to the large surface area of the nasal cavity:

This mucosal surface introduces additional acoustic energy loss in the low-frequency range. . . [L]osses due to viscosity, heat conduction, and wall impedance in an acoustic tube are all proportional to the ratio of the surface area to the cross-sectional area of the tube. Thus the bandwidths of the low-frequency poles and the zero for a nasal vowel are expected to be substantially greater. . . However, measured average bandwidths for the zero f_z and the additional pole F_n are about 200 Hz (Chen 1995). The introduction of nasalization appears to add about 100 to 200 Hz to the bandwidth of the first formant.

Stevens (1998: 316) comes to the following general conclusion about the spectral envelope of nasalized vowels: “[T]he calculated transfer functions for both front and back nasal vowels is that the spectrum shape at low frequencies (up to, say, 1200 Hz) is flatter and does not contain narrow or dominant spectral prominences.”

To review, the acoustic consequences of vowel nasalization, derived from over five decades of research, are these:

1. Widening of the bandwidth of F1 (and, for back vowels, F2);
2. Introduction of a pole-zero pair that prevents any one low-frequency resonance from being dominant; and
3. Introduction of another pole-zero below F1 due to acoustic coupling to a sinus, again preventing the dominance of one low-frequency spectral peak.¹⁷

¹⁷Stevens (1998: 306) observes that the coupling of the sinuses and the nasal cavities introduces “local fixed-frequency prominences in the spectrum as a consequence of additional pole-zero pairs in the transfer function of the combined vocal and nasal tract.”

Walker (2000: 69) observes, “It is well-known that nasalization tends to obscure the perceptibility of vowel height contrasts [F1], evidenced, for example, by the universal generalization that the number of vowels in a language never exceeds the number of oral vowels”¹⁸ (Ruhlen 1975, 1978, Bhat 1975, Crothers 1978, Beddor 1983, Wright 1986, Padgett 1997).

The acoustic consequences of nasalization are necessarily conditioned by the degree of velopharyngeal opening (i.e. coupling of the nasal and oral passages) as well as the total volume of the nasal passage. Because of the rather intricate structure of the paranasal sinuses and the inflammations and secretions that commonly block the ostia which connect the sinuses and the nasal passage proper, the geometry and volume of the entire nasal tract are difficult to calculate and are, moreover, highly variable among individual speakers (Kurowski and Blumstein 1993).

It has been suggested that the acoustics of nasalization can be ‘mimicked’ by other speech articulations, including voiceless fricatives (Ohala 1993, Ohala and Amadaor 1981). The wider-than-normal glottal opening that characterizes typical voiceless fricatives can result in some acoustic coupling with the sub-glottal cavity, resulting in increased bandwidth of F1 for adjacent vowels. Ohala (1993: 158) reports that “single period vowels excised from the portion of vowels immediately adjacent to voiceless fricative[s] and then iterated into 300–500 ms vowels were judged to be nasal by listeners.” Ohala (1993) cites phonological data which seem to be explained by this, e.g. spontaneous nasalization and nasal effacement—both phenomena transpiring near fricatives. The conclusion is that the glottal state during voiceless fricatives can spread to adjoining vowels and give an appreciable percept of nasalization.

Like the acoustics, the aerodynamics of the vocal tract are also substantially altered due to aperture of the velopharyngeal port. Here, the outcome seems much more straightforward: Where once there was relatively high pressure throughout the oral cavity and pharynx, the opening of the velopharyngeal orifice allows air an alternative escape route, thereby decreasing pressure throughout the system. Because both are conditioned by the same physical mechanism (i.e. the lowering of the soft palate), the acoustic modifications ascribed to nasalization are inextricably linked to the accompanying aerodynamic changes. Harkening back to Equations 1.5 and 1.6, one can easily see how a drop in pressure behind

¹⁸Of course, by itself, this does not explain why vowel height is impaired by nasalization. One could also have fewer nasal vowels than oral vowels if, for example, F2 distinctions were obscured.

the oral constriction (affected by the widening velopharyngeal orifice) negatively impacts both the volume velocity of the oral flow and the potential turbulence created at/near the oral constriction.

Because the velum is a relatively slow-moving articulator (Bell-Berti 1993, Krakow 1993, Moll and Daniloff 1971), it has long been observed that nasalization can occur epiphenomenally during segments that precede or follow nasal consonants. This is known simply as “nasal coarticulation.” Warren and Dubois (1964) use nasal flow evidence recorded along with the utterance *Are you home, papa?* to demonstrate the aerodynamic effects of the phenomenon.

In their experiment, Warren and Dubois (1964) observed the velopharyngeal orifice began to open for the /m/ in *home*. As early as the glottal fricative /h/, nasal flow was detected. Nasalization increased during the vowel and nasal consonant, then came to an abrupt halt during the closure of the first /p/ in *papa*. First we will consider the nasal onset. During the production of /h/, the area of the velopharyngeal aperture rose from a minimum area of 0 mm² to about 20 mm² in 125 ms. At this point, the voice onset of [ō] occurred. Thus, while air was flowing through the oral cavity to produce /h/ there was a relatively small degree of nasal flow during which the velopharyngeal orifice increased in size at a rate of approximately 0.16 mm²/ms. If we employ exact attention to articulatory detail, this 125 ms of frication should be transcribed as the nasalized glottal fricative [h̃].

We will now look at the nasal offset in this utterance. From a maximal velic opening of 80 mm², the value fell to 0 mm² over about 300 ms. Thus, while air was passing through the glottis to increase the pressure of the oral chamber, the velopharyngeal port was technically open, though constricting at a rate of -0.27mm²/ms. At the release of /p/ the velopharyngeal port was entirely closed, allowing no more air to escape through the nose and thereby satisfying the intra-oral pressure requirements characteristic of a voiceless stop.

Thus we see that even during the nominally oral consonants /h/ and /p/, articulatory evidence forces us to regard at least part of their production as nasal.¹⁹ While the nasally coarticulated /h/ may best be transcribed as [h̃], it does not seem quite right to transcribe the /p/ as [p̃] along the same principles, since the release burst of the plosive

¹⁹At least in perceptual terms, the position of the velum is irrelevant during /h/. There is not much evidence suggesting that it would sound different from a nasalized variant. Thus, allophonic variation of the type [h]~[h̃] may be widespread and practically unnoticeable. By regarding /h/ as an oral consonant by default, I simply follow conventional descriptions of the sound.

consonant is not expected to contain any nasalization (and in any case, the nasalization of the closure is adequately represented by the preceding /m/). For fricatives following nasals, there is also evidence of nasal coarticulation, reported as a ‘lag’ in nasal airflow extending into the fricative (Ali et al. 1979).

Moreover, Bell-Berti (1980) noted that velic lowering and raising bear a stable temporal relation to the achievement of the oral constriction. According to Stevens (1998: 43), “The minimum duration of an alternating movement of the soft palate that produces a single complete cycle from a closed velopharyngeal port to an open port and back to a closed port is estimated to be in the range of 200 to 300 ms.” Comparable durations were obtained by Krakow (1993) during fast speech. This is significant because it suggests that the movement of the velum is not necessarily conditioned by a fast or slow speech rate, the hypothesis being that in fast speech more segments adjacent to the nasal will become nasalized (Bell-Berti and Krakow 1991).

The extent to which a speaker can exert motor control over the velopharyngeal mechanism is still debated. Many early studies assumed that there was only a binary (open/closed) distinction for velic position. Bell-Berti (1993), however, argues for a more comprehensive view of soft palate position which includes intermediate states of opening. She notes “the problem of separating the intra- and intersegmental functions of the velum is further compounded by the almost constantly changing spatial relationships among the articulators” (1993: 64). The observation that velic height differs gradiently according to vowel height (usually the velum is low for low vowels and raised for high vowels) seems to indicate that intermediate positions of the velum are routinely used in language (Brücke 1856, Czermak 1869, Nusbaum et al. 1935, Moll and Shriner 1967, Moll 1962, Lubker 1968, Fritzell 1969, Bell-Berti et al. 1979, Henderson 1984). The perceptibility of different levels of nasalization, however, is a different matter. There is at least one language, Aceh (Malayic, Indonesia), that makes a phonemic distinction between “heavy” and “light” nasalization, but such a distinction is quite rare and may further imply the difficulty of controlling velic movements in any intermediate range (Durie 1985).

Physiologists have studied the internal composition of muscles controlling soft palate movement, finding in general that these muscles are not well-equipped to send much detailed information about movement and position to the brain. Muscle spindles are among the types of sensory receptors located in muscle that can provide information about proprioception and kinesthesia. They encode information primarily about muscle stretch. Before

research conducted by Liss (1990), spindles had been found exclusively in the *tensor veli palatini* and *palatoglossus* muscles (Lubker 1968, Lubker and May 1973, Lubker et al. 1972). Liss uncovered spindles in *levator veli palatini* (LVP) as well. Nonetheless, muscle spindles in LVP were relatively small and morphologically different from typical limb spindles or from spindles found in other speech mechanism musculature such as the jaw, larynx, lips, tongue, and respiratory system. Most of the evidence seems to indicate that a wide range of velic movements cannot be consciously controlled by speakers. This leads us to seriously entertain the conclusion that it is difficult to exercise precise control over the particular moments at which nasalization will start and stop during any given utterance.

Before leaving the aeroacoustics of nasals, it will be helpful to make a few observations about what is known to happen to oral obstruents that, by chance, become nasalized. The best source for such data is the literature dealing with velopharyngeal inadequacy. For example, Warren et al. (1993: 128–129) present evidence that cleft palate speakers actively compensate for the loss in resistance imposed by velopharyngeal impairment. They compare peak intraoral pressure in human subjects with three degrees of velopharyngeal inadequacy to the peak “intraoral pressure” in a passive mechanical system with dimensions matching those of an idealized human vocal tract. The result is striking, demonstrating that what is lost in terms of valvular pressure (when the velopharyngeal orifice opens) increases with greater output from the lungs. “[S]ubjects adopted active respiratory responses in an attempt to maintain pressure, and the strategies used were fairly successful in accomplishing this goal” (Warren et al. 1993: 131–132).

Finally, it is also worth noting that numerous studies have shown that when vowels are produced in the environment of nasal consonants, the position of the soft palate is lower for the low vowel /a/ than for the high vowel /i/ (Moll 1960). However, there is debate as to whether this should be considered a phonetic universal (Hajek 1997, Shosted 2006a). The possibility that this may be true of Brazilian Portuguese, Hindi, and/or French informs the choice of stimuli outlined in Section 2.5.

1.4 The Ohalian hypothesis considered

In the following sections, I will discuss five significant publications that have served to outline the Ohalian position on the status of nasalized fricatives (Ohala 1975, Ohala and Ohala 1993, Ohala et al. 1998, Yu 1999, Solé 1999). The relative merits and deficiencies of

the studies are addressed.

1.4.1 Ohala (1975)

Ohala (1975: 300) first argued against the existence of nasalized fricatives in general terms of the incompatibility of nasalization and oral obstruency:

Nasalization would be least compatible with oral obstruents. . . since the noise of fricatives and affricates and burst at the release of stops requires a build up of air pressure in the oral cavity. This would require that no air leak out of the oral cavity into the nasal cavity.

While Ohala admitted that it would be possible to produce voiceless fricatives like [s] with “some small velic leakage” he concluded that “it is extremely doubtful that voiced fricatives could be produced with a detectable amount of nasalization.” The author was aware of claims by Anderson (1975) regarding the existence of $[\tilde{v} \tilde{\delta}]$ ²⁰ but presumed that their acoustic realization must be similar to that of $[\tilde{w} \tilde{j}]$, i.e. frictionless continuants.

Ohala argued that fricatives, characterized by high oral pressure (vis-à-vis subglottal pressure) would be debilitated by velic opening. To maintain airflow through the glottis, it is necessary to maintain a sufficient pressure drop Δp with respect to the sub- and supra-laryngeal systems (the fluid mechanics of this phenomenon are discussed in Section 1.2). Specifically, pressure must be lower above the larynx than it is below the larynx in order for (egressive) speech to occur.²¹ With no supralaryngeal outlet (i.e. when the soft palate is raised and the mouth is closed), the air pressure above and below the larynx tends to stabilize and voicing eventually ceases. Voiced fricatives (like voiced stops) require lower oral pressure to maintain voicing and would be especially sensitive to a drop in pressure behind the oral constriction (Ohala 1983: 201–202). Thus, according to Ohala (1975), nasalized voiced fricatives are particularly untenable. One possible corollary of the argument as set forth is that voiceless fricatives are more resistant to small amounts of velic leakage because their oral pressure is higher than that of voiced fricatives.

It seems likely that pharyngeal and glottal fricatives (those articulated upstream of the soft palate) may be nasalized because nasal venting does not restrict fricative noise

²⁰Wondering what IPA symbols might be used to transcribe nasalized fricatives, Ohala (1975: 300) observed that “for $[\tilde{v}]$ IPA does recognize $[\text{ɱ}]$.” This usage of $[\text{ɱ}]$, the labiodental nasal ‘stop’, seems unsatisfactory because it fails to emphasize the oral flow that is characteristic of purported nasalized fricatives. Accordingly, I will use $[\tilde{v}]$ to symbolize the labiodental nasalized fricative, e.g. of Umbundu (see Sections 1.5.1, 1.7.9).

²¹The same is true, of course, for respiratory expiration.

generation. With respect to nasalized fricatives articulated upstream of the velopharyngeal port, Ohala (1975: 301) concluded that they are possible for two reasons: (1) velic opening would not prevent the build-up of air pressure behind a glottal or pharyngeal constriction; (2) “[N]oise produced by voiceless glottal and pharyngeal obstruents is so diffuse, so low in intensity, and with higher frequencies dominating in the spectrum that oral-nasal coupling would have little acoustic effect on it.” In other words, while pharyngeal and glottal nasalization are physiological possibilities, these are not likely to be adopted in any language due to problems with perceptibility.

1.4.2 Ohala and Ohala (1993)

Ohala and Ohala developed these ideas further in a 1993 paper. Previous conjecture on the incompatibility of obstruency and nasality was presented with the elocutionary force of a theorem (227):

Theorem 1.1 (Buccal obstruents require velic closure) *The velic valve must be closed (i.e., the soft palate must be elevated) for an obstruent articulated further forward than the point where the velic valve joins the nasal cavity and the oral cavity.*²²

The authors ascribe an aerodynamic “purpose” to the buccal constriction, i.e. to build up air pressure which, when released, will create audible turbulence. They remarked that failure to seal the nasal from the oral chamber would lead to leakage behind the constriction and through the nose, effectively reducing or perhaps eliminating entirely the requisite pressure drop across the oral constriction. This debilitated pressure-drop, they observe, is the hallmark of cleft palate speech.

Ohala and Ohala (1993) recognized that the existence of nasalized fricatives in any language would undercut the theorem. However, they were careful to note that the existence of such a fricative could only be substantiated through instrumental verification of velic position. Crucially, they noted that “one need not take the presence of nasalized vowels next to these sounds as unambiguous evidence” of nasalization during the fricative itself (1993: 228). They also cite a personal communication with Elmar Ternes (21 August 1991), in which the author of the influential study on Applecross Scots Gaelic (1989) (see Section 1.7.1) indicates that his claim regarding the existence of nasalized oral fricatives

²²The authors note that for them ‘buccal’ means “any place of articulation that is forward of the point where the velic valve joins the oral and nasal cavities” (Ohala and Ohala 1993: 227).

was based on “kinesthetic sensations during the imitation of these sounds.” Ternes himself reportedly agreed with the need to verify velic aperture during these purportedly nasalized sounds.

1.4.3 Ohala, Solé, and Ying (1998)

Based on previous conjecture regarding the status of nasalized fricatives, Ohala et al. (1998) approached the question experimentally. Two trained phoneticians (two of the study’s authors) uttered steady-state voiced and voiceless “strong” and “weak” fricatives. Pressure behind the oral constriction was bled intermittently through a tube of variable diameter (thus variable impedance²³) which had been inserted through the buccal sulcus of the speaker and behind the back molars. The tube thereby simulated velic leakage with variable pseudo-velopharyngeal vent cross-sectional areas. For the experiment, intraoral pressure was sampled using a catheter that had been directed into the pharynx through the nose. It was shown that changes in amplitude and quality of frication were related to the diameter of the pseudo-velopharyngeal vent. Specifically, a vent area of approximately 18 mm² decreased amplitude and fricative energy causing sibilants to sound more like nonsibilants. Furthermore, it was shown that for a given vent area, intraoral pressure was diminished less for voiceless than for voiced fricatives (following the hypothesis in Ohala (1975) (see Section 1.4.1). Presumably, the effect on the pressure drop across the constriction was weaker in voiceless fricatives because the open glottis in these segments allowed greater airflow from the lungs to compensate for the velopharyngeal loss. For the smallest catheter, 7.9 mm², pharyngeal pressure was not significantly affected. Moreover, there was no detectable effect on the quality of the fricatives under these conditions.

Ohala et al. (1998) found that a reduction in the magnitude of the pressure drop across the oral constriction caused voiced fricatives to become frictionless continuants. Furthermore, their results showed that aperiodic acoustic energy in the higher frequencies was reduced for voiceless fricatives. As Walker (2000: 67) notes, “The findings of this study clearly support the claim that nasalization is antagonistic to fricative sounds; however, this antagonism appears gradient such that the greater the velo-pharyngeal aperture, the greater the reduction in frication, and conversely, the smaller the velo-pharyngeal aperture,

²³Because the diameters of the tubes differed from the diameters of the actual velopharyngeal passage, the impedances of the two systems were not comparable. This limit on diameter was imposed by the fact that the tube had to be inserted behind the back molars.

the less perceptible the nasalization.” While Ohala et al. (1998: 3085) conclude that “the aerodynamic requirements for fricatives seem to be relatively narrow and unforgiving,” this study also indicates that fricatives may undergo a relatively minor degree of nasalization with little or no acoustic consequence. Scholars such as Walker (2000: 67) have concluded that nasalized fricatives “[D]o occur in some languages, although typically either degree of frication or perceptibility of nasalization will suffer in the production of these segments.”

1.4.4 Yu (1999)

In another experimental study, Yu (1999) investigated a diachronic phenomenon associated with the development of Mandarin from Middle Chinese. According to the author, high vowels in Mandarin assimilated in place of articulation and frication to preceding sibilants. He notes, however, that this assibilation pattern is systematically absent when the vowel is followed by a nasal consonant. Yu proposes first that Middle Chinese vowels articulated before nasal consonants were regressively nasalized. Second, he proposes that velic leakage during the articulation of such a contextually nasalized vowel was sufficient to sap pharyngeal pressure, oral volume velocity, and oral particle velocity (vis-à-vis that of high oral vowels). He hypothesizes that “when pharyngeal pressure is vented significantly during the opening of the velic valve, the necessary pressure build-up behind the constriction of a fricative is severely diminished, resulting in no audible turbulence” (1999: 341). This hypothesis is supported by an experimental investigation comparing pharyngeal pressure, volume velocity, and particle velocity for nasal and oral vowels in recorded utterances of an American English speaker.

The implications of Yu’s study for the status of nasalized fricatives is clear: nasalization reduces oral turbulence. The instrumental results suggest a principled, physical explanation for the absence of fricative vowels in nasalized contexts.²⁴ Frication cannot be produced in environments where velic leakage has bled pressure behind the oral constriction. Application of these results to the controversy of nasalized fricatives suggests that fricatives in the context of nasalization must suffer some loss of turbulence (the result of high particle velocity, as described in Section 1.2).

²⁴In a potentially related (synchronic) matter, Brazilian Portuguese word-final high nasal vowels [ĩ ũ] cannot devoice in word-final position, whereas their oral counterparts can. So, for example, [sapatu] ‘shoe’ is acceptable but *[atũ] ‘tuna’ is not. A similar state of affairs is reported to exist in Jivaro (Jivaroan, Ecuador) (Beasley and Pike 1957).

1.4.5 Solé (1999)

Solé investigated the role of aerodynamic factors in shaping phonological structure. Specifically, she discussed how aerodynamic factors, in combination with other constraints of production and perception, determine feature cooccurrence restrictions, i.e. why certain combinations or features in segments are likely to occur whereas others are rare or fail to occur.

This study emphasized the aerodynamic conditions required for trilling and frication in association with the features [VOICE] and [NASAL]. Solé analyzed the aeroacoustic effects on trills and fricatives caused by artificial variation of voicing and nasality. This was done through the instrumentality of a pseudo-pharyngeal valve that vented oral pressure (cf. Ohala et al. (1998), Section 1.4.3).

Intraoral pressure (P_o) was intermittently vented using catheters of varying cross-sectional areas (7.9, 17.8, 31.7, and 49.5 mm²), all inserted into the mouth via the buccal sulcus and the gap behind the back molars. Differences in catheter size were intended to simulate the effects of various degrees of velopharyngeal aperture. Audio and aerodynamic signals were recorded simultaneously under normal and artificially vented conditions. Subjects wore earphones through which white noise was played at a loudness sufficient to mask the high frequency noise of the fricatives. This was intended to discourage auditory feedback to the speaker, who upon hearing a debilitated fricative might compensate for the acoustic deficiency with an increase in subglottal pressure and hence transglottal flow.

It was found that velic openings less than or equal to 17.8 mm² did not significantly impair frication. The author concluded that such small velic apertures would be insufficient to create the percept of nasalization in adjacent vowels, so, too, on fricatives. This supposition is based on Maeda's 1993 designation of 40 mm² as the threshold for a "robust percept of nasalization on vowels." Warren et al. (1993: 143) lower this threshold considerably, however, concluding that listeners "usually perceive hypernasal resonance when the velopharyngeal opening for nonnasal consonants is greater than 0.10 cm² [10 mm²], and there is almost always some hypernasality perceived when the opening is greater than 0.2 cm² [20 mm²]" (see Section 1.3). If Warren et al.'s threshold is applied to Solé's results, then they could reasonably be construed as evidence *against* the Ohalian hypothesis, i.e. showing that frication is not adversely affected by a range of velopharyngeal apertures (10 ≤ 17.8 mm²) clinically shown to contribute a perception of "hypernasal resonance."

1.5 Against the Ohalian hypothesis

The most conspicuous challenges to the Ohalian view of nasalized fricatives were presented in studies of Umbundu (Schadeberg 1982) and Coatzospan Mixtec (Gerfen 1999, 2001). The phonetic patterns in these languages are reiterated in Sections 1.7.9 and 1.7.3, respectively. Because these two authors presented their results as rejoinders to the work of Ohala and his colleagues, in Sections 1.5.1 and 1.5.2 I consider how Schadeberg and Gerfen contextualized their results with respect to the Ohalian hypothesis. I also discuss potential weaknesses in their methodologies. Numerous additional languages reported to have nasalized fricatives are cited and described in Section 1.7, but they are not reviewed in the present section because of the authors' neutral stance with regard to the nasalized fricative controversy. While the description of Waffa (Stringer and Hotz 1973), for example, was apparently uninformed by the Ohalian hypothesis (in fact, it was published two years before Ohala (1975), so it could not be), the interpretation of data by Schadeberg and Gerfen seems directed at disproving the hypothesis. For this reason, I review the writings of these authors under special heading here.

1.5.1 Schadeberg (1982)

Schadeberg (1982) discussed a possible counterexample to the theorem stated in Ohala and Ohala (1993: 227). He claimed that Umbundu (Niger-Congo, Angola) in fact possesses a nasalized voiced fricative [ṽ]. However, as Ohala and Ohala point out, instrumental verification of air pressure build-up (i.e. obstruency) during the sound was not conducted by Schadeberg and has not been conducted, so far as I am aware, to this day. The challenge is to prove experimentally that the labiodental “fricative” reported by Schadeberg is not merely a nasalized glide [w̃] or, for that matter, a nasalized vowel [ū]. It is crucial in this case to find aspects of aperiodic, high frequency noise associated with fricative production, and to demonstrate that they are debilitated by nasalization (e.g. by comparison with oral [v]). A search for such acoustic cues of frication might be futile, however, since [v] is realized in many languages as a nearly frictionless approximant. Still, Schadeberg (1982) does posit a nasalized labiovelar glide [w̃] for the language, though it is doubtful that there are many (if any) minimal pairs contrasting the two sounds [w̃ ṽ].

Schadeberg (1982) also seemed to take exception to Ohala's (1975) observation that counterclaims were based only on a few South American and Celtic languages. Indeed,

it appears that Ohala was quite right about situating the nasalized fricative phenomenon geographically in South America, where it appears that most attestations do in fact occur (see Tables 1.8 and 1.9).

Schadeberg presents his reader with only four words in which [ṽ] occurs. This count is arrived at using a collection of approximately 2,000 lexical items gathered by the author, apparently in the field (no dictionary is cited).

In addition the lexical infrequency of [ṽ], there are a number of reasons why one might be sceptical of these findings. One of the four words exhibiting the questionable nasal fricative, óku-tyáãà ‘to cut firewood,’ was not in the dialect of Schadeberg’s three informants, who preferred [ɲ] to [ṽ] for this token. In a footnote Schadeberg (1982: 109) reports that “all the data on which this article was based” were provided by three female informants from Bié, who referred to the dialect of Huambo (which apparently none of them spoke) as “probably” having [ṽ] in the debatable word.²⁵ Third, the author notes that in another of the four words, ólu-néãa ‘reed,’ [ṽ] varies with [v]. Unfortunately, Schadeberg (1982: 118) “did not check whether nasalization is an optional possibility” in the three other words where stem-initial C1 is followed by non-nasal [v]. The author admits that “with so few examples [two to four], distributional restrictions and oppositions are difficult to establish.” Schadeberg (1982: 118) nonetheless observes that [v m mb] can all occur in the same positions as [ṽ] (as well as [ɲ]) and concludes that the nasalized fricative “has to be accepted as a rare but valid member of the phonological inventory of U[m]bundu.”

In Umbundu, nasalization occurs word-finally in monosyllabic stems, which consist of -CGV, -CV, -GV, or -V (G=glide), extending from the (final) nasal vowel over the entire word-final sequence “whenever phonetically possible” (Schadeberg 1982: 115). The fricative [v] nasalizes in word-final sequences (note that [s] cannot). The so-called ‘pure’ nasals [n m ɲ p] are never found in word-final monosyllabic stems.²⁶ No contrast exists “between nasalized and non-nasalized voiced continuants followed by [a nasal vowel]” because of leftward-spreading nasalization (Schadeberg 1982: 115). Most commonly, nasalization occurs in VCV sequences (with all the segments nasalized). If [v l j h w] appear before a nasal VCV sequence “it is difficult to decide whether these segments do or do not fall under the domain of nasalization” (Schadeberg 1982: 116). The author claims, against

²⁵Bié and Huambo are two central provinces of Angola that share a border approximately 200 miles long. Significant interaction between the inhabitants of the two regions could be expected.

²⁶There is no phonemic contrast between [n m ɲ p] and another nasalized continuant posited by Schadeberg, viz. [l̃], in this morphological context, though there are contrasts in other environments.

the judgments of his informants, that the nasalization in these consonants is weakly audible. Nasalization does not cross pre-stem boundaries, except weakly. Nasalization can be strongly realized on all the segments only if the first vowel is found in the stem, thus [óva-lá] or [óva-lá] (where, following Schadeberg’s convention, an under-tilde signifies weak nasalization *not* creakiness, as in modern standard IPA usage).

As one might expect, granting phonemic status to [ṽ] serves a broader phonological end. It is in fact helpful to Schadeberg’s analysis of nasal harmony in Umbundu, which works out more economically if nasal continuants are the locus of nasalization instead of nasal vowels (or nasal consonants themselves). Strangely, in Schadeberg’s analysis, vowels nasalize near nasal continuants but *not* next to ‘pure’ nasals like [n m ŋ ɲ]. According to the author, however, this is not strange at all (Schadeberg 1982: 127). His reasoning involves the “considerable articulatory effort” required to produce voiced nasalized continuants. “The nasalizing of adjacent vowels seems a natural consequence of this special effort, and it certainly helps the hearer to perceive the nasal quality of the obstruents” (Schadeberg 1982: 127). It is not clear whether by “articulatory effort” he refers to increased subglottal pressure (and therefore, transglottal flow), increased velopharyngeal opening, or both. The balance would certainly be a delicate one: Increased supraglottal pressure (a result of increased transglottal flow) would tend to extinguish voicing and increased velopharyngeal port size would tend to extinguish frication. But if it is nasalization that ‘spreads’ to segments adjacent to nasalized continuants like [ṽ] in Umbundu, then the widening of the velopharyngeal port is the only possible gesture to which Schadeberg (1982) could be referring. Assuming this to be the case, the “considerable articulatory effort” that goes into nasalizing the continuant must serve as the undoing of the fricative itself, venting back pressure more drastically with every incremental increase in aperture. On the other hand, if by “considerable articulatory effort,” the author referred to increased transglottal flow (to increase supraglottal pressure and thereby maintain oral frication in the face of an open velopharyngeal port), then one might expect partial voicing of adjacent segments, rather than nasalization, as the coarticulatory outcome.

In summary, the author unfortunately presented no instrumental evidence justifying his claim that the velum is lowered during the articulation of the labiodental nasalized fricative in Umbundu. He made reference to two phonetic degrees of nasalization, weak and strong (though the distinction is not phonemic as in Aceh (Malayic, Indonesia) (Durie 1985)), regrettably without aerodynamic or acoustic data to back up the proposition. The

author further argues for the existence of other nasalized continuants as well, viz. [ḥ ḷ ḷ̃ ṽ]. According to Schadeberg (1982: 110), [ḥ ḷ̃] are “relatively common,” [ḷ̃] much less so, and [ṽ] is “very rare.” In fact, the sound occurs in only about 0.02% of his lexical database. The case for increased frication in [ṽ] is complicated by the claim that only nasalized continuants (not the typical nasal consonants like [m n ŋ]) can cause coarticulatory nasalization in Umbundu. This could mean that the velopharyngeal port is opened wider for the nasalized fricative [ṽ] than it is for a consonant like [ŋ].²⁷ This increased aperture would be especially detrimental to a voiced fricative like [v] because a loss of back pressure (extinguishing frication) could be compensated only by increased subglottal pressure, which would critically imperil voicing.

If, on the other hand, Schadeberg’s data is taken at face value, it means that it is somehow possible to vent oral pressure (enough to create a percept of nasalization) and still generate perceptible orally-produced fricative noise. Such a state of affairs would present a strong challenge to traditional mechanical and aerodynamic models of the vocal tract.

1.5.2 Gerfen (1999, 2001)

Gerfen (2001) is an abbreviated version of Gerfen (1999: 121–211), a chapter on nasalization from his dissertation on the phonology of Coatzacoapan Mixtec. In the article, Gerfen (1999) sets his observations about nasalized fricatives in the context of a larger discussion regarding “what can constitute a speech sound in natural language” (Catford 1977, Lindblom 1990, Maddieson 1997, Ladefoged and Everett 1996). His data “challenge standard assumptions regarding the universal possibilities of nasalization,” viz., that buccal fricatives (especially the voiceless variety) should be incompatible with nasal venting. His thesis states, “It is the morphological nasalizing context which triggers anticipatory velum lowering in voiceless fricatives.”

Like Schadeberg (1982), Gerfen was aware of Ohala’s (1975, 1993) claim that substantial velopharyngeal aperture would siphon off the pressure build-up needed to create fricative noise across an oral constriction. He observes that Cohn’s (1993) survey of nasalized fricatives provides only a few possible counterexamples, including Umbundu (Schadeberg 1982), Waffa (Stringer and Hotz 1973), and Ìgbò (Carnochan 1948, Williamson 1969) (see Section 1.7). Commendably, Gerfen presents aerodynamic evidence to back up his counter-

²⁷The percept of nasalization on adjacent vowels could be caused by other factors, as well, including differences in spectral dynamics, length, etc.

claim, that the nasalized fricatives $[\tilde{s} \tilde{j} \tilde{\delta} \tilde{\beta}]$ exist in a Mixtec language of southern Mexico and must be accounted for in any set of phonetic universals.

1.5.3 Coatzospan overview

Coatzospan Mixtec (Mixtecan, Mexico) is a language that shows evidence of nasal harmony, i.e. the systematic propagation of nasal resonance from a specified start-point to a specified end-point within a word. The specification of these points (‘segments’ in more phonological terms) seems to vary widely across languages (Walker 2000). In Coatzospan Mixtec, formation of second-person familiar (2FAM) verbs involves the right-to-left propagation of nasalization from vowel to vowel. Intervening voiced consonants do not block the spread of nasality but voiceless consonants do. Gerfen’s provocative contention is that the very fricatives which stop the propagation of nasality (the voiceless consonants) can themselves be nasalized in the process. Thus, voiceless fricatives in Coatzospan Mixtec are TRANSPARENT (allowing the propagation of nasality) and MALLEABLE (able to undergo nasalization themselves) with respect to nasal harmony (see Table 1.7 for more about these terms. The term MALLEABLE in this context is unique, so far as I know, to my dissertation. It is not used by Gerfen (1999, 2001) or Walker (2000)).

The second person familiar (2FAM) of Coatzospan Mixtec verbs is formed by regressive nasalization within the domain of what is commonly called a ‘couplet’ in the Mixtecan tradition (either a CVCV or CVV syllable) (Pike 1948). Only the CVCV pattern is of interest here, since the medial C may be in some cases a nasalized fricative. This nasalization comes about under the effects of 2FAM nasal harmony, which involves the leftward propagation of nasality from vowel to vowel. If the medial consonant in CVCV syllables is voiced, then the leftmost vowel may be nasalized. Gerfen (2001) calls this a “transparent” consonant, though evidently he does not use this term in the same sense as Walker (2000), for whom TRANSPARENT denotes a consonant that may itself become nasalized (note that I call ‘nasalizable’ segments MALLEABLE; see Table 1.7).

Table 1.1: Fricatives through which nasalization ‘spreads’ in Coatzospan Mixtec. These fricatives are both TRANSPARENT (allow nasalization to ‘spread’) and MALLEABLE (become nasalized themselves).

BASE FORM		2FAM	
$\beta i \delta e$	‘wet’	$\beta i \tilde{\delta} \tilde{e}$	‘you (FAM) are wet’
$ku \beta i$	‘die’	$k u \tilde{\beta} i$	‘you (FAM) will die’

According to Gerfen (1999, 2001), voiceless medial consonants *do not* allow nasalization to ‘spread’ through (see Table 1.2 for examples).

Table 1.2: Fricatives that block nasalization in Coatzospan Mixtec. These fricatives may themselves be nasalized in the process, i.e. they are MALLEABLE. In any case, nasalization does not spread leftward, as in the tokens found in Table 1.1, e.g. *[kũtsĩ]. Note that [ũ ɪ̃] are non-modal creaky vowels.

BASE FORM		2FAM	
kũtsi	‘bathe’	kũtsĩ	‘you will bathe’
kĩfĩ	‘come’	kĩfĩ	‘you will come’

These segments that ‘block’ nasalization to an adjacent vowel may be MALLEABLE to nasalization, i.e. [s ʃ] may themselves become nasalized. This may give rise to some confusion, since voiceless fricatives in Coatzospan Mixtec are TRANSPARENT in Walker’s (2000) terminology but OPAQUE according to Gerfen (2001); these are not competing claims, but definitional ambiguities. As noted below (see Table 1.7), I have adopted the term MALLEABLE to describe segments that may be nasalized and NONMALLEABLE for segments that cannot be nasalized despite the ‘spread’ of nasalization ‘through’ the segment. The term TRANSPARENT refers generally to segments that allow nasal ‘spread’ (encompassing both MALLEABLE and NONMALLEABLE varieties) while the term OPAQUE refers to segments that disallow nasal ‘spread’ altogether (see Table 1.7 for a summary of these terms).

Gerfen (1999, 2001) presents aerodynamic evidence to claim that not only the voiced TRANSPARENT segments [β ð] may be phonetically nasalized but the voiceless OPAQUE segments [s ʃ] may be nasalized as well.

Gerfen’s instrumental approach (1999, 2001)

The author investigated the phonetic characteristics of segments that behaved as TRANSPARENT with respect to nasal harmony. Three female speakers participated in the study while in their home village of San Juan Coatzospan, Oaxaca, Mexico. A small foam plug known as a nasal olive was inserted in one of the speaker’s nostrils while the speaker manually plugged the other nostril. The pressure signal from the nasal olive was electrically transduced and recorded (Gerfen 1999: 14–18). Audio was simultaneously captured using a “close-talking” microphone worn by the speaker (it is presumed that the microphone was head-mounted). Unfortunately, the electrical output of the transducer was not calibrated at the time of the experiment, so the real-world values of nasal flow (e.g. in ml/s) at the

time of the experiment are unknown. A calibration of the transducer was performed later at the UCLA phonetics lab, so estimated values of nasal flow were later provided, but the standard error of this secondary calibration is unknown. This being the case, we have no idea how a calibration performed at San Juan Coatzospan might have differed from the calibration later performed at UCLA.

In an appendix, Gerfen (1999: 232–285) reproduces numerous diagrams of his aerodynamic data, indicating nasal flow during some fricatives and a lack of nasal flow during others. No systematic statistical analysis of these data is undertaken. The flow traces are presented anecdotally, i.e. as incidents whose variable occurrence remains unexplained. Moreover, the aerodynamic data is presented along with audio data in only one figure, and in this figure no calibrated scale of airflow has been provided (Gerfen 1999: 185, Figure 112). It is therefore impossible to tell from this study the effects nasalization might have on fricative acoustics. To be fair, this was not Gerfen’s research objective. It seems he intended to present anecdotal evidence of nasalization during some fricatives in Coatzospan Mixtec in order to construct a phonological model of the phenomenon. To the extent that we can rely on his methodology of data collection (including the unfortunate post hoc calibration of the instrument), we might say that he has been successful in this endeavor.

Recommendations

There are a number of problems with the methodology employed by Gerfen (1999, 2001). By outlining them here, I hope to show how the methodology employed in the present study (see Chapter 2) may fill in some of the gaps.

First, it is important to remember that what is typically measured in airflow studies is air *pressure* behind some sort of resistance (Cotes et al. 2006: 61–62). The pressure drop, Δp between two arbitrary points in the flow can be approximated from the Navier-Stokes equation:

$$\Delta p = p_a + p_c + p_f \tag{1.7}$$

where p_a is the pressure increment or decrement due to linear acceleration between the two points, p_c is the pressure change due to convective acceleration between the two points, and p_f is the pressure change due to frictional losses. In a pneumotachograph, p_a and p_c are both minimized by the design of the instrument: the former by placing the pressure ports close together and the latter by ensuring that the inlet and outlet diameters are equivalent.

In this manner, it can be said that $\Delta p = p_f$. By referring back to Equation 1.2 (Poiseuille’s Law) we can substitute p_f for Δp

$$\Delta p = p_f = U_l(8\eta\ell/\pi r^4) \quad (1.8)$$

where (to review), η is the gas viscosity, ℓ is the length of the tube, and r is the radius. With ℓ and r controlled in the design of the pneumotachograph, it turns out that the pressure drop is linearly related to flow and is dependent on gas viscosity. The linear relation between the pressure drop and flow is crucial, since it is ultimately flow, not pressure, which we would like to extrapolate from the analysis.

In Gerfen (1999, 2001) pressure is also measured, using a so-called ‘nasal olive’. As with a pneumotachograph, the standard assumption in using this device is that the pressure drop across the device will relate in a linear way to nasal flow. However, to obtain a signal of sufficient strength, it was necessary for Gerfen’s subjects to *close one of their nostrils*, thus increasing the pressure build-up in the nasal cavity. The condition of the second closed nostril does not obtain during normal speech, so it may be argued that Gerfen’s data are compromised by the methodology. With one nostril open, the air pressure was presumably not robust enough to be measured accurately.²⁸

For Gerfen, it was also unfortunately necessary to calibrate the nasal flow device after returning from the field. As Gerfen himself notes, the reported nasal flow rates are only a “rough approximation.” His results may in fact bear little relation to the actual values. Moreover, the transducer system was calibrated at a single flow rate, viz. 250 ml/s. Multiple flow rates (at least three) are needed to demonstrate the crucial presence of a linear relationship between the physical input to the transducers and the electrical output. Without basing a calibration on at least three flow rates (even once the author had returned from the field), it is entirely possible that the transducers behaved in a non-linear fashion. If that were the case, this would further compromise his data.

Under these conditions, demonstrating that nasalization levels are roughly comparable to those during a nearby nasalized vowel is the next best solution. While this seems convincing in some cases, it is unclear how great the difference is in others. A statistical analysis, minimally accessing the ratio of peak nasal flow during the fricative and peak nasal

²⁸Gerfen’s methodological compromise may in fact be taken as supportive of the Ohalian hypothesis, i.e. that with a certain degree of velopharyngeal leakage, the buccal pressure will not be sufficient to produce a fricative. In other words, the weakened nasal pressure signal produced with the second nostril open is analogous to the weakened oral pressure signal that would be generated with the velopharyngeal port open.

flow during the subsequent nasal vowel would have gone a long way to clarify the matter.

As mentioned previously, aerodynamic measures were gathered through a nasal olive. Since the nasal olive was inserted in only one nostril, the speakers had to plug the other one manually in order to prevent leakage. Gerfen (2001) refers to an objection raised by John Ohala, viz. if one of the nostrils is occluded, “the spiking present during the production of these fricatives may simply be an artifact of slight velum raising (but not opening) which compresses the air trapped in the nasal cavity between the velum and the nostrils.” Gerfen (2001) addresses this concern with four arguments in favor of his hypothesis:

1. “It is highly unlikely that this amount of air could be moved by a slight raising gesture of the velum when it is already in a position to seal the velopharyngeal port”;
2. “Nasal flow is sustained in a number of tokens” and a one-time raising gesture anticipates a spike, not continuous nasal flow;
3. The nasal flow trace should trend negative toward the end of the fricative as the velum begins to lower in preparation for the release into a nasalized vowel;
4. Measurements indicate “obviously” that the velum is in a lowered position during the first vowel and at the onset of the fricative.

I will mention several ways in which these defenses are unsatisfactory. First, any “amount of air” referred to by the author is unquantifiable due to the calibration difficulty discussed earlier. Second, the author makes no attempt to define the notion of “spike” versus “continuous nasal flow” or quantify it in relation to the number/kind of tokens in which such phenomena occur. It seems disingenuous to state the spiking is atypical of the data when even a cursory glance at the nasal flow diagrams provided in the appendix show that “spikes” in nasal flow are highly characteristic of voiceless nasalized segments like [š ʃ] and what might reasonably be called “continuous nasal flow” is characteristic of the voiced fricatives like [ð ß] (Gerfen 1999: 232–285). In any case, without a quantitative analysis which utilizes some mathematical definition of “spike” versus “continuous” flow, this is merely an argument, as it were, “in the eye of the beholder.” Third, negative nasal flow is most difficult to substantiate without an accurate calibration and/or ‘landmarks’ in the signal where nasal flow is known to be zero (e.g. during oral stop closures) (Shosted and Willgoos 2006). In summary, Gerfen (2001) was unable to invalidate Ohala’s point,

especially in the absence of a scientific assessment of when “spikes” do and do not occur in his data. The simplest course of action in resolving the matter would be to perform a nasal olive experiment in the laboratory. Under more controlled conditions, the range of airflow discontinuities produced by raising the soft palate could be determined.

It would have been advantageous to his analysis had Gerfen also recorded the fricative oral flow. Decreased oral volume velocity during the nasalized fricative (*vis-à-vis*) an oral fricative would have helped to substantiate the reallocation of transglottal flow through the nasal chamber. To see a decrease in oral flow that accompanies the nasal spikes would provide crucial reassurance to the sceptic.

Gerfen’s anecdotal observations of the acoustic signals (reported to the reader in little detail) indicate that Coatzospan Mixtec nasalized fricatives are not “frictionless continuants” as Ohala and Ohala (1993) reason $[\tilde{\beta} \tilde{\delta}]$ must be. From the single figure provided, it appears the fricative is fairly noisy (Gerfen 1999: 185, Figure 112). Strikingly, however, there is no scale provided for the nasal flow in this figure, so it is virtually impossible to correlate the actual degree of nasalization with any change in the acoustic signal. It is not clear that there is any change in fricative amplitude associated with nasalization, but due to the lack of (even an imperfect) calibration scale, it is impossible to tell how nasalized the fricative is in the first place.

Gerfen concludes that nasal fricatives are indeed infelicitous segments, since “velum lowering has negative aerodynamic and acoustic consequences for obstruency” (Gerfen 2001). This seems an odd claim to make after failing to demonstrate (or argue) that nasalization of fricatives has an appreciable effect on the acoustics of the fricatives themselves. Moreover, he makes no attempt to assess their oral flow characteristics. It would seem more natural for Gerfen to conclude that velum lowering does *not* make any significant difference, at least among Coatzospan Mixtec fricatives. While it does not seem unreasonable that nasal flow should exist during the production of a fricative sound (cf. (Solé 1999), Gerfen (1999, 2001) did not rigorously assess the relationship between aerodynamic and acoustic variables for the Coatzospan Mixtec nasalized fricatives. Hence, the significance of his data remains unclear.

1.6 Strong and weak versions of the hypothesis

Especially in relation to Gerfen’s work, it may prove helpful to differentiate a strong and a weak version of the Ohalian hypothesis concerning nasalized fricatives. The strong version is stated in Hypothesis 1.1:

Hypothesis 1.1 (Strong version) *Nasalized fricatives cannot exist phonetically.*

This version may be derived from early postulatory writings such as Ohala (1975, 1983). A weaker version, based on the empirical studies of Ohala et al. (1998), Solé (1999), Yu (1999) might read like Hypothesis 1.2. Corollary 1.1, an addendum to the weak version of the hypothesis, has gone unstated in the literature yet seems like a natural extension thereof. The production side of this corollary will be the main focus of the present study. Assessments of Hypothesis 1.2 and Corollary 1.1 are presented with respect to the findings of the present study (Chapter 3) in Chapter 4.

Hypothesis 1.2 (Weak version) *Nasalized fricatives, if they exist, must be acoustically debilitated.*

Corollary 1.1 *Due to their acoustic debilitation, nasalized fricatives are not phonologized in any language.*

Unyielding pursuit of the strong hypothesis could have some undesirable consequences. For example, what would one make of the fact that cleft palate speakers routinely produce nasalized sounds that are also orally fricated (though certainly to variable degrees) (Weinberg and Horii 1975)? An awareness of research on cleft palate speech is evident in, e.g. Ohala and Ohala (1993), but studies of cleft palate fricatives in particular are not addressed. Ohala takes no position on nasalized fricatives in cleft palate speech; in effect, he does not deny that such fricatives exist.

Thus, Gerfen’s aerodynamic evidence in favor of nasalized fricatives demonstrates the untenability of Hypothesis 1.1 (the strong version) but remains silent on Hypothesis 1.2 (the weak version) and Corollary 1.1. Regarding the matter of phonologization (Corollary 1.1), Schadeberg (1982: 127) approaches the subject by mentioning the “considerable articulatory effort” expended in the production of [ṽ]. This is of course an imprecise and unsatisfactory statement in scientific terms, but it is at the very least a vague intimation

of why nasalized fricatives are not commonly phonologized in the world’s languages. Gerfen (1999, 2001) unfortunately does not address the matter of phonologization, though his aeroacoustic data might have been used for this purpose. Had Gerfen (2001) shown no statistically significant relationship between nasal flow and frication intensity, his results would have discredited the weak version of the hypothesis as well. As it stands, Gerfen is not in a position to refute the arguments of Ohala and Ohala (1993), Solé (1999), Yu (1999) because his data say nothing significant about the reduction in spectral energy that may (or may not) be the hallmark of a nasalized fricative. Gerfen’s results suggest that nasal leakage sometimes occurs when fricatives are produced adjacent to nasalized vowels. The acoustic consequences of this phenomenon still await discussion.

1.7 Reports of nasalized fricatives

The following sections include data relating to nasalized fricatives in a typologically and geographically diverse set of the world’s language. The list is exhaustive, according to my own knowledge and that of various sources, particularly Cohn (1993) and Walker (2000). Most of these reports were not originally presented in relation to the Ohalian hypothesis but as mere descriptions of the phonological inventories and/or grammars of the languages at hand (excepting Coatzospan Mixtec (Gerfen 1999, 2001) and Umbundu (Schadeberg 1982)).

1.7.1 Applecross Scots Gaelic (Celtic, Scotland)

As of 2001, there were 183 residents of Applecross, Ross Shire, Scotland and at that time only 31.2% or approximately 60 individuals could “speak, read, or write” Gaelic (Highland Council 2004).

Ternes (1989) presents a phonological analysis of nasalization in the Applecross dialect of Scots Gaelic. Among other things, Ternes’s study is known for positing a number of voiceless nasalized fricatives. He argues that instead of attributing phonemic vowel nasalization to vowel segments, it should be attributed to consonants instead. For example, he claims that [t^hã:v] *tàmh* ‘rest, repose’ is underlyingly and historically /t^h a:ĩ/. He also posits such forms as /sa.ĩhux/ [gloss not provided] and /k^hĩxk/ [gloss not provided]. Ternes’ main argument in positing these nasal fricatives seems to be one of elegance or economy of analysis, claiming that establishing only a few nasal consonant phonemes “would be limited and would certainly not exceed the number of nasalized vowels and diphthongs required” for

competing interpretations (Ternes 1989: 132). He mentions two problems for his phonological account, neither of which touch on the aerodynamic implausibility of anterior nasalized fricatives. Interestingly, one of the problems deals with forms where there are no consonants which he considers “susceptible” to nasalization—only stops which are in his words “per definitionem excluded from nasalization.” From an aerodynamic standpoint, it has been argued that fricatives are also unsusceptible to nasalization, e.g. (Ohala and Ohala 1993, Ohala 1975). Ternes winds up rejecting the nasal consonant analysis, but not because aerodynamic fricatives are a problem. In fact, he again posits them in an alternative analysis, that of the “long nasal component” (Ternes 1989: 133).

Ternes justifies this final alternative in this manner:

“By not having to decide whether phonemic nasality should be attributed to consonants or to vowels, the drawbacks inherent in either solution are avoided, while at the same time their respective advantages are accumulated” (Ternes 1989: 133).

The analysis is comprised of the following constraints:

1. The center of nasalization lies in the vocalic nucleus of the stressed syllable of a stem. Nasalization is strongest in the center. From the center, nasalization extends in a forward and backward direction unless or until checked by a further condition;
2. In the backward direction, nasalization comprises the consonantal onset of the stressed syllable, but never extends beyond;
3. In the forward direction, nasalization may extend as far as the end of the word, unless checked by (4) or (5);
4. Nasality does not extend beyond stops;
5. The vowel phonemes /e o ə/ never function as the center of nasalization. The nasal ‘long component’ (which the author argues should not be termed a ‘nasal prosody’), obeys no constraints with respect to fricatives *per se*. As long as the fricative occurs in a place relative to the nasal vowel that is not checked by constraints (2-5), it is, in the author’s estimation, nasalized.

Ternes posits a number of phonetic forms that seem implausible from an aerodynamic point of view (Table 1.3:

Table 1.3: Nasalized fricatives of Applecross Scots Gaelic

ANTERIOR			POSTERIOR		
š	t ^h āhũšk	‘senseless person, fool’	š̃	kānāš̃	‘sand’
š̃	š̃āhũk	‘axe, hatchet’	h̃	š̃ōh̃i	‘tame’
ř̃	ř̃ēnēvāř̃	‘grandmother’	ř̃	strāř̃iř̃	‘string’
f̃	f̃riāṽ	‘roots’			
č̃	āhũč̃	‘neck’			

So-called “vibrants” (the author does not indicate whether these are multiple-strike articulations) are also supposedly affected by the long nasal component, e.g. [māh̃h̃āř̃] ‘mother’, [ř̃ṽūāř̃] ‘to dig’. If the author is referring to nasalized trills, and if his findings are valid, it would represent a direct counterexample to the instrumental work on nasalized trills conducted by Solé (1999) (see Section 1.4.5 for a discussion of Solé’s study).

1.7.2 Chichimeco-Jonaz (Otopamean, Mexico)

Lastra (1984) mentions only one nasalized fricative in Chichimeco-Jonaz, an Otopamean language of Guanajuato state, Mexico. In 1993, the language was spoken by 200 individuals in San Luís de la Paz, Jonáz village (Gordon 2005). The sound of interest is nominally a nasalized, voiced labiodental fricative [ṽ]. However, Lastra observes (p.c. 2006) that there may be little or no contact between the teeth and upper lip during its articulation. Younger speakers of Chichimeco-Jonaz (unsurprisingly) tend to replace [ṽ] with Spanish [β]. For this reason, it may be quite difficult to document the acoustic and aerodynamic specifications of the sound, even in the proximate future.

1.7.3 Coatzospan Mixtec (Mixtecan, Mexico)

An Oto-Manguean language of northern Oaxaca, Mexico, Coatzospan Mixtec is spoken by about 5,000 individuals (500 monolinguals) in the village of San Juan Coatzospan (Gordon 2005). According to Gerfen (1999, 2001), speakers of Coatzospan Mixtec routinely nasalize fricative segments that occur adjacent to nasal vowels. Gerfen (1999) presents nasal flow evidence (gathered with a nasal olive) suggesting that the velum is substantially lowered during the production of the erstwhile oral fricatives [f̃ ð̃ ṽ] when these adjoin a nasal vowel. Crucially, Gerfen (1999, 2001) does *not* argue that the fricatives of Coatzospan Mixtec are phonemically nasalized. However, he clearly argues that nasalization does coöcur with oral

frication. The details of fricative nasalization in Coatzospan Mixtec are comprehensively reviewed in Section 1.5.2.

1.7.4 Epena Pedee (Choko, Colombia)

Harms (1994, 1985) asserts that fricatives may be nasalized in Epena Pedee, a Choko language spoken by approximately 3,500 people on the Pacific coasts of Colombia. According to Harms (1994: 8), “Nasalization is a suprasegmental feature that is associated with the syllable and spreads to the right within a word.” Moreover, “Any segment within a nasal syllable (whether derived or inherently nasal) is manifested in the form of its nasalized variant.” Epena Pedee has the phonemic fricatives /s h/, but [ɸ ɣ ʁ β] occur allophonically in word-medial position. Harms mentions nothing that would preclude the nasalization of these segments as well, and one example of a nasalized bilabial fricative, [náβ̃ẽ] ‘mother’, is in fact cited. Other nasalized fricatives occur in [sĩẽĩĩ] ‘sugar cane’ and [wãĩĩⁿdá] ‘go.PAST’.

1.7.5 Ìgbò (Niger-Congo, Nigeria)

Ìgbò is a language of Nigeria reported to have five nasalized fricative phonemes, including [h̃] (Williamson 1969: 87). The putative alveolar nasalized fricatives [š̃ ž̃] undergo palatalization before [i], resulting in two more nasalized fricatives at the surface level, [ʃ̃ ʒ̃].

Williamson (1969: 91) observes that nasalization “runs through the entire syllable” in Ìgbò. According to Cohn (1993: 332), this makes the analysis of the Ìgbò nasalized fricatives “less problematic” than if the nasalized fricatives were purely phonemic. Nonetheless, Williamson (1969: 87) cites a number of disyllabic words that seem to have only one underlying nasal segment, thus making it unclear how the distinction may be considered non-phonemic (Table 1.4).

Table 1.4: Nasal and oral fricatives in Ìgbò

ORAL			NASAL	
s	isà	‘to spread out’	ĩsà	‘to wash (face/pot)’
ʃ	àfĩ	‘bead’	ĩfĩ	‘six’
z	izù	week (of four days)	ĩzù	‘to steal’
ʒ	oʒi	‘message/errand’	eʒĩ	‘pig’

While Ladefoged and Maddieson (1996: 132) accept Carnochan’s (1948) docu-

mentation of [h̃] in Central Ìgbò, they are more sceptical of Green and Igwe’s (1963) report of nasalized voiced and voiceless labiodental and alveolar fricatives. Rather than having simultaneous nasal and oral airflow, these segments are probably oral fricatives that occur with nasalization of the following vowel—“the device of marking the consonants as nasalized being employed, as noted by Williamson (1969), to identify the limited set of consonants that can begin syllables with nasalized vowels” (Ladefoged and Maddieson 1996: 132).

1.7.6 Icelandic

Icelandic has a relatively large speaker population (240,000), compared to other languages that reportedly have nasalized fricatives (Gordon 2005). Walker (2000: 65) explains that descriptions of Icelandic “are explicit in claiming that nasal airflow is maintained during the fricative,” citing Pétursson (1973) and Einarsson (1940).

Pétursson believes that *constrictives nasales* (nasal continuants) exist in Icelandic. He describes the formation of these sounds as a relaxation of consonantal stricture when a nasal precedes a homorganic continuant (“Devant des constrictives homorganes les occlusives relâchent leur articulation et deviennent des constrictives”) (1973: 116). However, he notes that there is considerable disagreement on the matter, citing Einarsson (1940), Poirot (1924), and Bergsveinsson (1941), all of whom have fundamentally different views.

Using kymographic recordings, Einarsson (1940: 462) argues that the nasal continuants have the same oral articulation as the following consonant:

If an *n*, at the end of a first element in a compound, or at the end of a word of a sentence, comes to stand before a spirant or a liquid except *h*, it usually loses the stop-formation and is turned into a homorganic nasalized spirant or liquid. These sounds are voiced, and the position of the organs seems to be the same as that of the following spirant or liquid, perhaps a bit more open.

This suggests that nasals occurring before fricatives are at least partially realized as voiced nasalized fricatives. However, Einarsson (1940: 463) observes that “there is no way of drawing the line where the [nasalized] vowel ends and the voiced spirant begins.” With the observation that cymograph recordings cannot settle the question unequivocally, Einarsson (1940: 464) determines that “nasalized spirants... are still so determined by the... auditory senses.” Unfortunately, the collection of an auditory impression does not by itself constitute a falsifiable experiment, a state of affairs that seems clear to Einarsson.

Pétursson (1973) is idiosyncratic in his transcription of the *constrictives nasales*, partially following Einarsson (1940). Pétursson uses subscript fricatives (always voiced) for nasal continuants preceding [s z θ ʃ] (e.g. [dan_zsa] for *dansa* ‘to dance’) and [ɱ] before the labiodentals [v f]. For the sake of consistency, I use standard transcriptions like [ž ɥ] to present the data in Table 1.5.

Table 1.5: *Constrictives nasales* in Icelandic, after Pétursson (1973)

FRIC	Ortho	IPA	
z	<i>dansa</i>	tanžsa	‘danser’
v	<i>umfram</i>	ɱɰv̥fram	‘en outre’
ð	<i>ennþá</i>	enðθau	‘encore’
j	<i>án hjarta</i>	aɲjçarta	‘sans coeur’
ɣ	<i>Svanhvít</i>	svaɲɣ̥x ^w it	personal name

There is no question that the major portion of the fricatives in these Icelandic words is articulated without nasalization, but there is some supposition that at least part of the fricative is produced with a significant degree of nasalization, and moreover, that this portion is voiced. However, there is no indication that the distinction between nasal continuants and occlusive nasals is phonemic. In fact, some of the examples cited by Pétursson (1973) arise only at word boundaries.

In opposition to the views of Pétursson and Einarsson, Bergsveinsson (1941) argues that nasals before fricatives are simply deleted, leaving residual nasalization on the preceding vowel. Poirot (1924) argues that the vowel undergoes compensatory lengthening and the nasal is realized with its original duration (“la voyelle aurait subi un allongement compensatoire et la nasale conserverait la moitié de sa durée normale”). Phonologists and phoneticians, therefore, differ substantially on how the Icelandic nasal continuants are actually realized (if at all). Though he does not use the term ‘fricative’, it is evident from his transcription and description of the sounds that Einarsson (1940) believed nasalized fricatives were relatively common phenomena in Icelandic speech.

1.7.7 Inor (Semitic, Ethiopia)

Inor (sometimes referred to by its Amharic designation, Ennemor or Ennämor) is an Semitic language of Ethiopia spoken by approximately 280,000 individuals (Gordon 2005). Though the language has a full range of fricatives, including [f f^w s z ʃ ʒ x x^w xʃ],

only [β] and [ʒ] are said to undergo nasalization (Hetzron and Marcos 1966). Chamora and Hetzron (2000: 10) observe that nasal harmony invokes the change [β] → [ɱ]. However, they do *not* claim that [β] is a fricative in Inor, rather that it is as an approximant. Chamora and Hetzron (2000) make no mention of the voiced alveopalatal nasalized fricative [ʒ̃] cited by Walker (2000). Since [β] is considered an approximant *before nasalization occurs* and [ʒ̃] is unsupported in the more recent analysis, there seems to be no compelling reason to keep Inor on the list of languages that purportedly possess nasalized fricatives.

1.7.8 Japanese

In Japanese, the syllable-final nasal has a number of allophones which range from a nasalized vowel to a nasal consonant homorganic with the following stop. In isolation the sound may be articulated as a “voiced frictionless nasalized prevelar spirant” (Bloch 1950: 102). Vance (1987) correctly observes that “frictionless spirant” is something of a contradiction in terms. It seems clear that this “debuccalized” or “underspecified” nasal segment is best described as a nasal velar approximant, perhaps resembling [ŋ] in its acoustic properties (Trigo 1988, Padgett 1991). The positing of a Japanese nasalized velar fricative [ɣ̃], as in Applecross Scots Gaelic, seems unwarranted and quite possibly unintended by Bloch (1950).

1.7.9 Umbundu (Niger-Congo, Angola)

Schadeberg (1982: 117) argues for the existence of the nasalized fricative [ɣ̃] in four words of Umbundu, a Bantu language spoken by approximately 4 million Angolans. Schadeberg (1982: 127) reasons that “considerable articulatory effort is needed to produce voiced nasalized continuants, much more than for the production of pure nasals” and this is precisely why he claims that nasal continuants [ṽ ḥ̃ ḷ̃ ȷ̃ Ẃ̃] are the locus of spreading nasalization—not nasal vowels and not the so-called ‘pure’ nasal consonants [n m ŋ ɲ] themselves. A fuller description of Schadeberg’s methodology, along with the presentation of his views with regard to those of Ohala (1975), are given in Section 1.5.1.

1.7.10 Waffa (Papuan, Papua New Guinea)

Waffa is spoken by approximately 1,300 individuals in Morobe Province, Papua New Guinea, at the headwaters of the Waffa river (Gordon 2005). Stringer and Hotz

(1973) indicate that Waffa has a voiced bilabial nasal fricative [β̃] which contrasts with [β m mb]. Table 1.6 illustrates words employing these segments in initial and medial positions (Ladefoged and Maddieson 1996: 134).

Table 1.6: Nasal contrasts in Waffa

INITIAL			MEDIAL	
mb	mbúumə	‘stamens’	símbáu	‘fly’
β	βíndi	‘man’	kóoβə	‘father’
β̃	β̃oóka	‘back, leech’	β̃aβ̃ə	‘skin’
m	mókoo	‘live coals’	β̃aimáura	‘tree’

1.7.11 Other ‘nasal harmonic’ languages

Walker (2000: 3) defines ‘nasal harmony’ as a phenomenon that “comes about when an underlyingly nasal segment, such as a phonemic nasal stop or nasal vowel, triggers the nasalization of an adjacent string of segments in a predictable and phonologized way.” Here, we are particularly concerned with Walker’s discussion of languages that allow ‘nasal harmony’ or ‘nasal spreading’ to cross fricative segments. Under the assumption that the velum is lowered when nasalization ‘spreads’ from one segment to another, cases in which a fricative intervenes between the ‘trigger’ and ‘target’ of nasalization may imply the existence of a nasalized fricative.

That a language allows nasalization to ‘spread through’ certain segments, however, does not necessarily entail that those segments are thereby nasalized. Indeed, Walker (2000: 61) differentiates between segments that allow spreading nasalization but do not themselves undergo nasalization (she calls these TRANSPARENT segments) and those that allow spreading nasalization but remain oral (she does not assign a term to these). For the sake of clarity, I will not follow Walker’s terminological choices.

In keeping with the harmony literature (and those terms that seem most clear for present purposes), I will refer to segments that allow the spread of nasalization as TRANSPARENT.²⁹ Those that block the spread of nasalization I will call OPAQUE.³⁰ To differentiate the two types of TRANSPARENT segments, those that may become nasalized

²⁹Walker calls these THROUGH segments.

³⁰Though I regret the divergence from Walker’s text, I feel that it will ease the comprehension of my own arguments.

and those that may not, I will use the terms MALLEABLE and NONMALLEABLE. These definitions are summarized in Table 1.7.

Table 1.7: Nasal harmony definitions

TERM	DEFINITION
TRANSPARENT	Allows nasalization to ‘spread’ through, e.g. rightward
OPAQUE	Prevents nasalization from ‘spreading’ through
MALLEABLE	Becomes nasalized when nasalization ‘spreads’ through
NONMALLEABLE	Remains oral when nasalization ‘spreads’ through

transparent fricative languages

As suggested in Table 1.7, all MALLEABLE and NONMALLEABLE segments must also be TRANSPARENT segments, otherwise their susceptibility to nasalization would remain unknown. Unfortunately, the grammars from which Walker drew her typological data do not consistently clarify whether the TRANSPARENT segments are MALLEABLE or NONMALLEABLE. Accordingly, I present the languages in Tables 1.8 and 1.9 as cases of *potentially* nasalized fricatives, i.e. TRANSPARENT fricatives (unless, of course, the details of an individual language, e.g. Coatzospan Mixtec, were discussed in an earlier section). Ideally, aeroacoustic analysis of all of these languages should be undertaken. Guaraní and Umbundu, which probably have the largest numbers of speakers, seem like good places to start (as Walker (2000: 242) notes in the case of the former).

Walker (2000) cites four languages in which vowels, glottals, glides, liquids, and fricatives are TRANSPARENT segments, whereas obstruent stops are OPAQUE segments (Table 1.8). This is the least common pattern in her nasal harmony database (Type IV in her typology) (Walker 2000: 65). According to her summary of the typological data, “This suggests that if the demand of nasal harmony is strong enough to spread through fricatives, it generally is strong enough to target some stops as well” (Walker 2000: 65).

Walker (2000: 64–65) cites 28 languages in which all classes of segments (vowels, glottals, glides, liquids, fricatives, and obstruent stops) are TRANSPARENT segments (see Table 1.9). These are called Type V languages in Walker’s typology. In Table 1.9, I have listed all of the fricatives that occur in each language, though in only a few cases have explicit claims been made about their status as MALLEABLE or NONMALLEABLE segments

Table 1.8: Type IV nasal harmony languages (Walker 2000). All segments in these languages, excepting obstruent stops but including fricatives, allow nasal harmony to ‘spread’ (i.e. they are TRANSPARENT segments). For Inor, (Chamora and Hetzron 2000) use the symbol for a bilabial fricative but categorize the (oral) sound as an approximant (the IPA symbol for a voiced bilabial approximant is [β̞]). Its nasalized counterpart is symbolized as [m̃], which the authors use to symbolize a labial (not labiodental) sound. Thus β̞ is given here in parentheses.

LANGUAGE	DIALECT	FAMILY	LOCATION	FRICATIVES
Inor		Semitic	Ethiopia	(β̞) ʒ
Epena Pedee		Choco	Colombia	s h
Itsekeri		Niger-Congo	Nigeria	ɣ
Scottish Gaelic	Applecross	Celtic	Scotland	f s ç ʃ x h
Umbundu		Niger-Congo	Angola	v h

(e.g. Coatzospan Mixtec (Gerfen 1999, 2001), Inor (Ennemor) (Hetzron and Marcos 1966),³¹ and Epena Pedee (Harms 1985)). Presumably, Walker does not include Icelandic among the 29 (though she specifically mentions Pétursson’s (1973) and Einarsson’s (1940) reports of nasalized fricatives in that language) because Icelandic phonology does not show signs of nasal harmony (Walker 2000: 65).

Walker (2000: 67) makes several typological observations regarding her database. “In the class of obstruents it is always the case that voiced fricatives are the most compatible with nasalization and voiceless stops are the least compatible. Continuancy and voicing thus are qualities favoring nasalization of obstruents. For segments with just one of these qualities, languages appear to vary in whether continuancy or voicing is more compatible with nasalization.” From her survey, it is clear that all languages which treat some obstruents as TRANSPARENT universally treat voiced fricatives as TRANSPARENT, but voiceless fricatives and voiced stops may sometimes trade places in the hierarchy. For example, for Applecross Scots Gaelic, voiceless fricatives are TRANSPARENT and voiced stops are OPAQUE but for Epena Pedee (Choco, Panama), Orejon (Tucanoan, Peru), and Parintintin (Tupí-Guaraní, Brazil), voiced stops are TRANSPARENT and voiceless fricatives are OPAQUE. At least for this sample, the second pattern seems to be more common, i.e. Walker (2000) cites only one language in which voiceless fricatives, but not voiced stops, behave as TRANSPARENT

³¹In a much later publication (and posthumously in the case of the second author), Chamora and Hetzron (2000: 17) eliminate [ʒ] altogether and indicate that /β̞/ is realized as [m̃] (which they confusingly refer to as a *labial*, not labiodental, approximant) under the effects of nasal harmony. More to the point, they categorize /β̞/ as an approximant. For these reasons, Inor should no longer be included among languages that purportedly have nasalized fricatives.

segments with respect to nasal harmony.

opaque fricative languages

While I have given considerable descriptive emphasis to those languages in which nasal harmony is allowed by fricative segments, this may appear to give undeserved statistical importance to such languages. In fact, according to Walker's typology, in a majority of nasal harmony languages, fricatives are OPAQUE to the spread of nasalization.

Walker's study includes a sample of 85 nasal harmony languages. Of these, 61% (n=52) block spreading nasalization. Though still appreciable, only 39% (n=33) are languages in which fricatives allow nasal harmony to pass through. Languages with fricative 'blockers' are typologically and geographically diverse, with a full range of fricatives represented (Walker 2000: 61-63). She cites Midwestern English, South Castilian Spanish, Silacayoapan Mixtec (Mixtecan, Mexico), Marathi, and Kolokuma Ijo (Kwa, Nigeria), among others, as languages in which fricatives prevent the regular activity of nasal harmony.

1.8 Summary

This review of the controversy surrounding nasalized fricatives has demonstrated a number of points:

1. Fricatives and nasals have antagonistic aerodynamic requirements: fricatives require high back pressure and nasals deplete it;
2. It is possible that different kinds of fricatives (employing different kinds of aerodynamic regimes) will be more or less affected by nasalization (e.g. the voiced vs. voiceless distinction has been mentioned (Ohala 1975), but the sibilant vs. non-sibilant distinction may also be of interest).
3. Based on aerodynamic/mechanical models of the vocal tract, the phonetic existence of nasalized fricatives has been questioned (the strong version of the Ohalian hypothesis) (Ohala 1975, 1983);
4. It has been postulated that fricatives, once nasalized, must lose some characteristic acoustic quality (the weak version of the Ohalian hypothesis) (Ohala and Ohala 1993, Solé 1999, Yu 1999);

5. It remains to be determined whether these acoustic characteristics are perceptually significant enough to explain why nasalized fricatives are rarely, if ever, phonologized in the languages of the world;
6. Despite the influence of the Ohalian hypothesis (or in some cases in response to it), nasalized fricatives have been explicitly reported in a number of geographically and typologically diverse languages. In only a single case (Coatzospan Mixtec) have reports of such fricatives been accompanied by recorded evidence of nasalization (Gerfen 1999, 2001).
7. Nasalized fricatives potentially exist in a much larger number of languages (many of them under-described) with nasal harmony (Walker 2000). Any language in which nasalization ‘spreads through’ fricative segments is potentially significant in this regard.
8. Most languages that experience nasal harmony do not allow nasalization to ‘spread through’ fricative segments.

In light of aerodynamic evidence suggesting the presence of nasalization during Coatzospan Mixtec fricatives and with the numerous accounts of nasalized fricatives in other languages (see Section 1.7), it is incumbent upon us to abandon Hypothesis 1.1 (the strong version) in favor of Hypothesis 1.2 (the weak version) and Corollary 1.1. The task, then, is to measure the effects of nasalization on oral frication. The methodology and outcomes of such an investigation will constitute the remainder of this thesis.

Table 1.9: Type V nasal harmony languages (Walker 2000: 64–65). All segments in these languages, including fricatives, allow nasal harmony to ‘spread through’, i.e. they are (TRANSPARENT). It is not known, however, whether the fricatives become nasalized in the process (i.e. whether they are MALLEABLE). An ‘*’ indicates that the whole inventory could not be determined and/or has not been reported.

LANGUAGE	DIALECT	FAMILY	LOCATION	FRICATIVES
Apinayé		Ge	Brazil	s z v ʒ
Barasano	Northern	Tucanoan	Colombia	s h
Barasano	Southern	Tucanoan	Colombia	s h
Bribri		Chibchan	Costa Rica	s z ʃ h
Cabécar	Southern	Chibchan	Costa Rica	f s ʃ x
Cabécar	Northern	Chibchan	Costa Rica	f s ʃ x
Cayuvava		(isolate)	Bolivia	β s ʃ h
Cubeo		Tucanoan	Colombia	v ð h
Desano		Tucanoan	Colombia, Brazil	s*
Epera		Choco	Panama	f s h
Gbeya		Niger-Congo	Central African Republic	s z f v h
Gokana		Niger-Congo	Nigeria	f v s z ʒ
Guanano		Tucanoan	Colombia	s h
Guaraní		Tupí	Paraguay, Brazil, Colombia	s ʃ x h v ɣ ɣ ^w
Guaymí		Chibchan	Panama	s x
Igbo	Ohuhu	Niger-Congo	Nigeria	f v s z ɣ h h ^w
Içuã Tupí		Tupí-Guaraní	Brazil	h
Kaiwá		Tupí-Guaraní	Brazil	v s ʃ h
Mixtec	Atatlahuca	Mixtecan	Mexico	*
Mixtec	Coatzacoapan	Mixtecan	Mexico	β ð ð ^j s ʃ x
Mixtec	Ocoatepec	Mixtecan	Mexico	β ð s z ʃ ʒ h
Orejon		Tucanoan	Peru	β s ʃ h
Parintintin		Tupí-Guaraní	Brazil	β h
Shiriana		Shirianian	Venezuela, Brazil	(Φ) s ʃ h
Siriano		Tucanoan	Colombia, Brazil	*
Tatuyo		Tucanoan	Colombia	h
Tucano		Tucanoan	Colombia	s h
Tuyuca		Tucanoan	Colombia, Brazil	*

Chapter 2

Method

2.1 Research hypotheses

Several hypotheses will be tested in the present study. They are extensions of Hypothesis 1.2, the weak version of the Ohalian hypothesis regarding nasalized fricatives, i.e. “Nasalized fricatives, if they exist, must be acoustically debilitated.”

1. Some acoustic qualities of fricatives are modulated by the presence of nasalization, or in mechanical terms, the opening of a vent behind the smallest constriction in the system;
2. These modulations increase as the degree of nasalization increases, or in mechanical terms, as the vent opening enlarges;
3. The acoustic modulation(s) associated with nasalized fricatives in human speech is/are comparable to the acoustic modulation(s) associated with mechanical nasalized fricatives produced by a vocal tract model (the design of which will be specified in Section 2.7.1);

It is not entirely clear in the nasalized fricative literature what these so-called ‘acoustic modulations’ might be. In this study, the following variables will be scrutinized under nasalized and non-nasalized conditions:

1. High-frequency energy (Shadle 1985, Stevens 1998, Solé 1999);
2. Spectral peak bandwidth (Johnson 1997, Stevens 1998);

3. Low-frequency energy (Delattre 1954, House and Stevens 1956, Hattori et al. 1958, Fant 1970, Fujimura and Lindqvist 1971, Bell-Berti and Baer 1983, Hawkins and Stevens 1985, Bognar and Fujisaki 1986, Fujimura 1962);

Traditional aeroacoustic models of the vocal tract suggest that high-frequency energy should be higher, spectral peak bandwidth should be lower, and low-frequency energy should be higher for oral fricatives vis-à-vis their nasalized counterparts (see Sections 1.2 and 1.3). However, with the exception of the hypothesis dealing with high-frequency energy (Solé 1999), none of these hypotheses has been verified for fricatives under the effects of nasalization.

2.2 Methodological overview

The research hypotheses in Section 2.1 will be verified using data from two different sources, viz. sounds produced by human vocal tracts (I will refer to these throughout as ‘spoken’ fricatives) and sounds produced by a mechanical model (‘mechanical’ or ‘model’ fricatives). Though there are various drawbacks in the acquisition and analysis of each type of data, it is hoped that when used in conjunction with one another they will increase our understanding of nasalized fricatives, if in fact they occur in human language.

The acoustics of each type of fricative (spoken and mechanical) will be assessed using the same techniques, including spectral analysis. Due to differences in the human and mechanical vocal tracts, the aerodynamics of each kind of fricative will be assessed in different ways, but the critical aerodynamic information will be recorded in each case. Thus, it can be said that the following constitutes an ‘aeroacoustic’ analysis, as it attempts to draw correspondences between the aerodynamic and acoustic features of the sounds involved.

Detailed information about each aspect of the methodology is given in this chapter. For convenience and clarity, however, the following brief summary is provided.

2.2.1 Spoken fricatives

Speakers produced voiceless fricatives under varying nasal conditions. Stimuli were VCV utterances where V1 and V2 were variably nasal and oral (both vowels had the same specification in this regard) and C was a buccal fricative (e.g. [ũfũ ufu]). Following the results of Ali et al. (1979), the presumption was that in some cases the fricatives or

portions thereof (especially the edges) would be nasalized. Nasalization was verified using a conventional oral and nasal air mask design. Thus, the acoustics of fricatives that were appreciably nasalized could be analyzed with respect to the research hypotheses of Section 2.1. Oral measures are also reported to substantiate the reallocation of transglottal flow through the nasal vent (as observed in the recommendations for improving the methodology in Gerfen (1999, 2001), Section 1.5.3). It must be noted that airflow is only an incidental indication of velic opening, but is commonly used in place of more direct (and necessarily invasive) measures (Cohn 1993).

2.2.2 Mechanical fricatives

Because the size of the velic opening during nasalized fricatives can only be measured indirectly (still using non-invasive means) and because the aerodynamic mask design for the organic fricatives precluded high quality acoustic recordings, a mechanical model of the post-velopharyngeal region of the vocal tract was constructed (see Section 2.7.1). The alveolar fricative [s] was modeled using articulatory data taken from an MRI study of American English fricatives (Narayanan et al. 1995). The size of the velopharyngeal vent was manipulated mechanically in order to produce the fricative under increasingly ‘nasalized’ conditions, i.e. by increasing the size of the vent diameter incrementally.

2.3 Languages

Spoken data were gathered from languages that have phonemically nasal vowels. Hindi, Brazilian Portuguese, and French have ten, five, and three such vowels, respectively. For the purposes of the present study only the so-called ‘corner’ vowels of each language were used. For Hindi and Brazilian Portuguese, the set includes [ĩ ũ ã], and for French [ẽ ã õ].

Though each language has voiced fricatives, the present study is limited to voiceless fricatives only. This decision was made for two reasons:

1. The model vocal tract constructed for this study did not allow for the production of ‘voiced’ sounds, so no comparison of spoken and mechanical fricatives could be made;
2. Because of their high air flow requirements, nasalized voiceless fricatives seem more controversial than nasalized voiced fricatives, at least with regard to the aerodynamic

hypotheses discussed in Section 1.6.¹

The voiceless buccal fricatives of Hindi are [f s ʃ]; for French and Portuguese they are [f s ʃ] and sometimes [ɸ], depending on the speaker. This last consonant may be realized as [x h χ] in Brazilian Portuguese and as a uvular or apical trill [ʀ r] in French.²

According to Ohala (1991), stress is not distinctive in Hindi; there is in fact controversy as to whether lexical stress even exists in the language. In French, stress is often described as falling on the last syllable of the word, except in connected speech (Fougeron and Smith 1999: 80). In Brazilian Portuguese, lexical stress typically falls on the penultimate syllable, but can occur in other positions; orthographically, these cases are signaled by a variety of diacritic markings.

2.4 Speakers

Three speakers of Hindi, two of French, and one of Brazilian Portuguese participated in the study. Two of the Hindi speakers were male, both from Delhi. The third Hindi speaker was a female, who reported that her parents were from Delhi but traces of Calcutta Hindi could also be found in her speech. The French speakers were both females, one from Paris the other from Normandy. The Brazilian Portuguese speaker was a female from Brasília. All speakers were UC Berkeley students, between 25 and 35 years old.

2.5 Stimuli

Speakers of each language uttered nonsense VCV syllables, where C was a buccal fricative (e.g. French [afɛ afɔ afa]). As mentioned previously, V was limited to a set of three, at the corners of the language’s vowel space, i.e. [ā ī ū] for Hindi and Brazilian Portuguese, [ā ē ɔ̃] for French. The syllables were composed of all language-appropriate sequences of V1, buccal fricative, and V2 in two nasal control groups. These groups consisted of different nasalization environments where either both vowels were nasal (VC̃Ṽ) or oral (VCV).

V2 was stressed in all tokens (to the extent that this is possible in Hindi; see Ohala (1991)). For example, the Brazilian Portuguese and Hindi speakers uttered the following

¹This is not to suggest that voiced nasalized fricatives should be accepted without further investigation. Nevertheless, it seemed prudent to constrain the scope of the present study.

²In the present study, the Brazilian Portuguese subject produced [x] in the nonsense syllables provided. The French speaker generally produced a non-fricative, which was therefore not analyzed.

stimuli, among many others: [ãsĩ' asĩ].

The intervocalic consonant was limited to each language's voiceless fricatives anterior to the velopharyngeal orifice (i.e., the so-called 'buccal' fricatives). The total number of stimuli for each speaker was therefore:

1. Hindi: $3 \text{ vowels} \times 3 \text{ fricatives} \times 3 \text{ vowels} \times 2 \text{ nasal control groups} = 54$;
2. Brazilian Portuguese: $3 \times 4 \times 3 \times 2 = 72$ (assuming the realization of /r/ as [x]); and
3. French: $3 \times 3 \times 3 \times 2 = 54$

Stimuli for each language were presented in native orthography, i.e. in Devanagari for Hindi and in the Roman alphabet for Brazilian Portuguese and French speakers. Since nonsense words were used and speakers were not trained to read the International Phonetic Alphabet, special consideration was given to the orthographic representation of the vowels and fricatives among the stimuli.

The Devanagari script provides a unique symbol for each sound, sometimes involving a combination of base symbol and diacritic marking(s) placed above and/or below this radical. Each fricative is represented by a unique base symbol in the script. Each nasal vowel is represented by drawing a dot above the corresponding oral vowel character. Vowels are represented through the use of diacritics when following consonants and as full characters when preceding them, but this presents no special challenge here.

In Portuguese, the low nasal vowel is represented through the addition of a tilde, e.g. *sã* [sã] 'healthy.FEM'; all other nasal vowels are represented by the addition of a following *-m* in word final position or before labials and *-n* elsewhere, e.g. *aipim* [aipĩ] 'sp. of cassava'; *onça* [õsɐ] 'jaguar'; and *ombro* [õbru] 'shoulder'. In Brazilian Portuguese, the grapheme *-s* is pronounced [z] in intervocalic position. Voiceless [s] can also occur in that position, but it is represented by *-ç*. The grapheme *-rr-* was used to represent the uvular/velar fricative. Word-final stress is typical when the final vowel is nasal or underlyingly high front, i.e. not /e/ raised to [i]). When stimuli contained a word-final oral vowel, stress (associated with unreduced vowel quality) was signaled through the use of standard Portuguese diacritics: *-ê* for word-final [e] and *-ô* for word-final [o].

In French, nasal vowels are represented using various orthographic strategies: word-finally, we observe *-in*, *-ain*, *-en* for [ɛ̃]; *-ent* and *-ant* for [ɑ̃]; and *-on* for [ɔ̃]. Word-initially, we observe *ain-* for [ɛ̃]; *an-* for [ɑ̃]; *on-* and *om-* for [ɔ̃]. According to convention,

the digraph *-ss-* was used to represent intervocalic [s] in French.

For Hindi, writing the stimuli in a form that could be understood by the subjects was no great challenge because of the direct symbol-to-sound correspondence in the Devanagari script. For Portuguese, where the situation is slightly more complicated, care was taken to use *-n* or *-m* as appropriate before consonants (V1), *-m* as appropriate in word-final position (V2), and standard diacritics for word-final stress. In French, where there were a number of orthographic possibilities for each nasal vowel, stimuli were analogized based on words like *ainsi* [ɛ̃si] ‘like this’; *pain* [pɛ̃] ‘bread’; *onze* [ɔ̃z] ‘eleven’; *saumon* [samɔ̃] ‘salmon’; and *antan* [ãntã] ‘yesteryear’. Accordingly, French [ɛ̃] was represented by *ain*; [ɔ̃] was represented by *on*; and [ã] was represented by *an*.

All stimuli were presented and exemplified to the speakers, using analogy to real words if confusion arose, during a short interview conducted before the recording sessions.

2.6 Spoken data

Simultaneous audio, nasal, and oral airflow signals were recorded for aerodynamic analysis. Later, a separate audio recording was made for acoustic analysis. For both sessions, speakers uttered the stimuli in the following frame sentences:

1. Brazilian Portuguese: *diz __ duas vezes* [dʒiz __ duɐʃ vɛziʃ] ‘s/he says X two times’;
2. French: *d’ __ descendit* [d __ dɛsɑ̃di] ‘s/he came down from X’; and
3. Hindi: [ʃəbd __ dek^h rəhə hɛ] ‘he is seeing the word __’.

Frame sentences were designed foremost to exercise prosodic control over each of the stimuli for a given language and to situate each utterance in an easy-to-define aerodynamic / acoustic context for later signal processing. In addition, the apical consonants ([z] and [d] in Brazilian Portuguese; [d] in French and Hindi) that surrounded the stimuli controlled the external-edge vowel transitions (i.e., the initiation of V1 and the terminus of V2). These transitions were not anticipated to have any particular consequence in the current analysis; nonetheless, in order to reduce the risk of introducing confounding variables, it seemed prudent to control for the effects of coarticulation in this manner. A recording of the sequence [ãfã] (Hindi) is given in Figure 2.1.

Audio, oral flow, and nasal flow were sampled simultaneously, as described in Sections 2.6.1, 2.6.2, 2.6.3. After the aerodynamic recording session, a separate audio

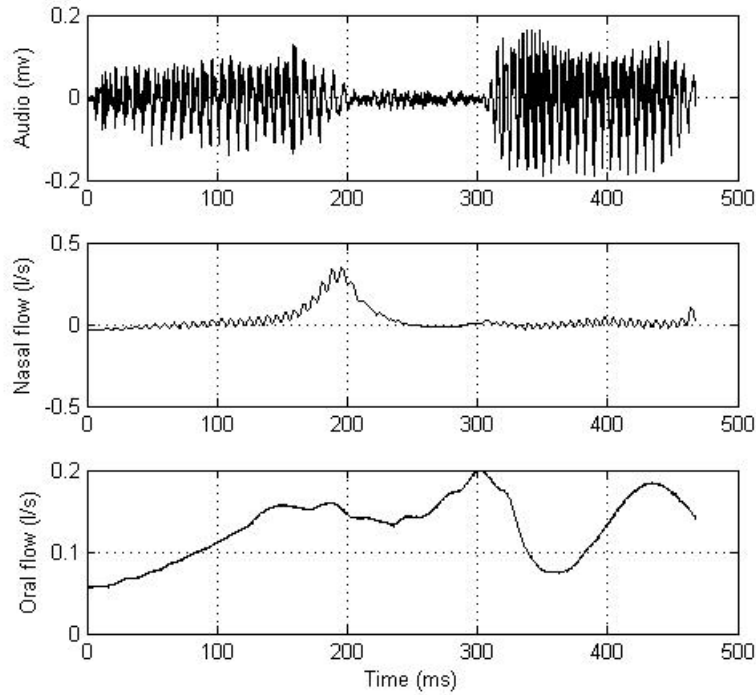


Figure 2.1: Audio, oral flow, and nasal flow recordings of the token [ãfã] (Hindi).

recording was made under conditions more appropriate to acoustic analysis, i.e. while the subject was *not* wearing oral and nasal masks.

All data signals were digitized at 20 kHz using a Dell Optiplex GX270 computer, a multifunction data acquisition board (Model PCI-6013, National Instruments Corp., Austin, TX) with a shielded connector block (Model BNC-2110, National Instruments Corp.), and Matlab 7.0.4 software running on a Windows XP platform in the Phonology Laboratory at the University of California, Berkeley.

2.6.1 Audio

For the aerodynamic session, audio was recorded using a cardioid dynamic microphone (frequency range 30 to 16,000 Hz) (Model D-190E, AKG Acoustics, Nashville, TN) positioned approximately 5 cm from the speaker's mouth and a dual microphone pre-amplifier (Model SX202, Symetrix, Inc., Mount Lake Terrace, WA). The audio quality was degraded by the oral mask, as described in Section 2.6.2. However, the audio signal was

still adequate for segmentation of the simultaneously-recorded aerodynamic signals. To overcome the problem created by the mask, audio was recorded a second time using a head-mounted microphone (Model SM10A, Shure Inc., Evanston, IL) and a Marantz solid state recorder (Model PMD670, D&M Professional, Itasca, IL) in a soundproof audiometric booth. For acoustic measurements (other than the segmentation of the aerodynamic signals themselves) all audio data comes from this second, higher-quality audio recording session.³

2.6.2 Oral flow

An oral mask (Model OM-2, Scicon R&D, Inc., Encino, CA) (Rothenberg 1977) was connected to a low-frequency transducer (model PTL-1, Glottal Enterprises, Inc., Syracuse, NY) via a length of tubing 10 cm long with an interior diameter of 0.5 cm. The output from the transducer was low-pass filtered (4-pole, Butterworth) at 75 Hz using an analog filter (Model 3364, Krohn-Hite Corp., Brockton, MA). The oral mask was held in place by the subject, who was instructed to maintain a snug fit, confirming that a seal was formed, in particular, at the upper lip and chin. The experimenter periodically verified the fit through visual inspection, especially during the production of low vowels, where jaw movement may cause slippage.

One critical drawback of an aerodynamic methodology that uses such masks is that the mask acts as a filter of the simultaneous acoustic signal.⁴ The amplitude of the sound pressure signal decreases (especially in the higher frequencies) when the mask is worn, but more importantly, the spectrum of the sound is altered significantly. Figures 2.2 and 2.3 illustrate these differences. Not only does the mask reduce the amplitude of the spectral frequency peak, it also introduces at least two spurious low frequency formants, presumably based on the geometry of the mask itself. In a study concerned with small amplitude changes in various frequency ranges, this is a significant problem.

Because the quality of the audio signals were compromised in this manner, it was necessary to make acoustic recordings unfiltered by the oral mask. This was accomplished

³For the second session, it was not possible to separate noise produced at mouth from frication at the nostrils (presumably the nasal mask was able to eliminate this noise during the aerodynamic session). Future experiments should contemplate ways to adequately prevent the conflation of the two without compromising the recorded data. Any friction generated at the nares in the contexts discussed would probably be small, but this has not been verified.

⁴While various methods were contemplated to get around this problem, including inserting a small microphone *inside* the mask, none of them has produced satisfactory results. For example, when the microphone is placed inside the mask, sound reverberation off the proximate walls of the mask produce a virtually unusable acoustic signal.

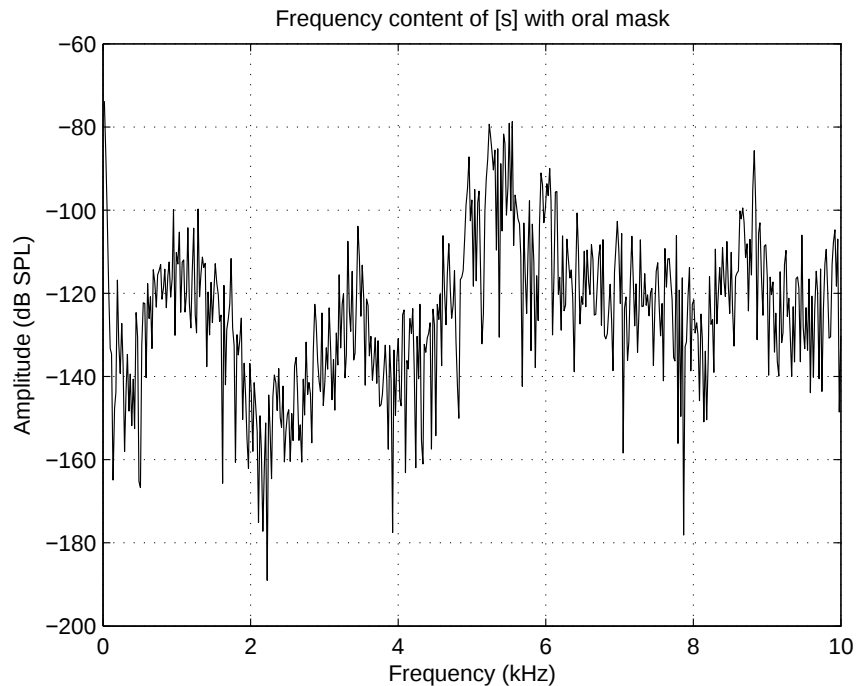


Figure 2.2: FFT of an alveolar [s] produced with speaker wearing Scicon OM-2 (oral mask).

using a dynamic head-worn microphone (Model SM10A, Shure Inc., Evanston, IL) in an anechoic chamber at the UC Berkeley Phonology Lab. Recordings were digitized to a Marantz solid state recorder (Model PMD670, D&M Professional, Itasca, IL). Unfortunately, using this methodology the acoustic signals of specific utterances could not be compared directly to their accompanying oral flow signals. Thus, the nasal and audio flow evidence is only generally indicative of the conditions that obtained during the ‘unfiltered’ recordings.

It was assumed that if nasal airflow during the fricatives in nasal syllables could be established as significant (with respect to fricatives in oral syllables), then the same effect should hold for recordings when aerodynamic records could not be made. While this arrangement is less than ideal, the constraints are imposed by the experimental instruments available. The methodology involving mechanical fricatives (Section 2.7.2) was conceived, in part, to compensate for this deficiency.

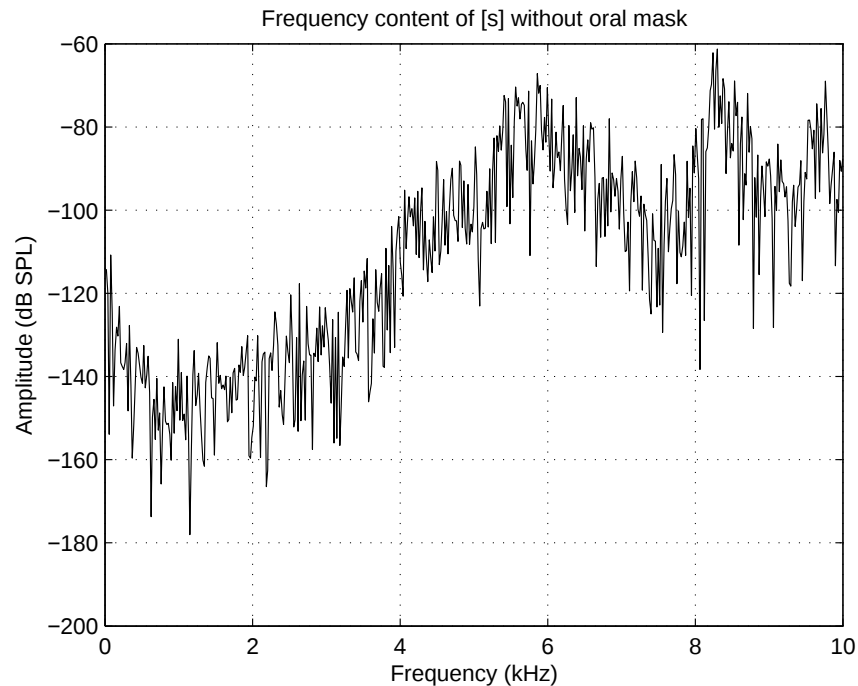


Figure 2.3: FFT of an alveolar [s] produced without the Scicon OM-2 oral mask.

2.6.3 Nasal flow

A nasal mask (GoldSeal model, Respironics, Inc., Murrysville, PA), intended for use in the treatment of adult obstructive sleep apnea (Brown et al. 1995), respiratory failure, and respiratory insufficiency, was used to sample nasal flow. The nasal mask was vented through its exhaust port using a piece of fine synthetic mesh and was connected to a wide-band transducer (model PTW-1, Glottal Enterprises, Syracuse, NY) via a length of tubing 10 cm long with an interior diameter of 0.5 cm. The output from the transducer was low-pass filtered (4-pole, Butterworth) at 75 Hz using an analog filter (Model 3364, Krohn-Hite Corp., Brockton, MA). The GoldSeal mask cushion is filled with gel which allows the mask to form a complete seal against the face.

2.6.4 Flow calibration

Procedure

A pneumotach calibration unit (Model MCU-4, Glottal Enterprises, Inc., Syracuse, NY) was used to calibrate the aerodynamic signals. This micro-processor controlled ‘artificial lung’ provides calibration sequences with user-selectable flow rates and flow volumes. Plaster negatives of the oral and nasal masks were fabricated by hand and used as mask gaskets. These were mounted on the vent of the calibration unit. This ensured that the oral and nasal masks fit snugly against the apparatus, increasing the chances that all of the vented air would be channeled towards the transducers. Airflow was expelled from the machine at five different flow rates, viz., -1000 cm³/s, -500 cm³/s, (0 cm³/s), 500 cm³/s, and 1000 cm³/s (at 1000 cm³ total volume). These values were related to the electrical responses of the PTL-1 and PTW-1 transducers using least-squares linear regression. Calibrations were performed before each speaker was recorded. It was hoped that repetition of the calibration procedure would yield increased accuracy, e.g. in the event of performance variations in the transducers between sessions.

The relationship between the electrical responses of the transducers and the known input of the calibration unit varied across languages, speakers, and tokens. It is not entirely clear why this should be the case, but the effect is probably due to small fluctuations in the behavior of the transducers, ambient temperature changes, and/or changes in the seal between gasket and mask. To determine the reliability of each calibration, the correlation coefficient for each calibration was calculated, as discussed below.

Correlation coefficient

The correlation coefficient (r^2) of the predicted versus actual responses of the measuring device is defined as:

$$1 - \frac{\sum_{i=1}^n (Y_i - Y_{p_i})^2}{\sum_{i=1}^n (\bar{Y}_i - \bar{Y}_{p_i})^2} \quad (2.1)$$

where Y is the actual measured response of the device and Y_p is the predicted response according to a least-squares linear regression model. For example, a correlation coefficient of 0.98 suggests that the linear fit used for a given calibration explains 98% of the variation

in the measured responses of the transducer. Accordingly, a higher correlation coefficient is indicative of a better fit and therefore a more reliable calibration.

If the correlation coefficient for a session was less than 0.95, the calibration was performed again. Fortunately, this occurred on only a few occasions, so it is assumed that the calibrations for the various sessions were reliable.

It should be noted, however, that a reliable calibration does not guarantee the accuracy of the results. Once the subject secures the mask, any slippage can reduce the accuracy of the aerodynamic recording, whether or not the transducers have been calibrated accurately. For this reason, it was necessary for the experimenter to pay attention to the seal of the mask (particularly the oral mask) around the subject's face. If the mask slipped in any observable way, the recording of the token was repeated. Other fluctuations in the response of the transducers, due to ambient temperature and/or humidity, were not controlled with any degree of precision.

2.7 Mechanical fricatives

2.7.1 Model design

The model was built of clear, removable acrylic plates (0.625 cm thick) drilled through with holes of various areas (ranging from 0.18 cm² to 7.92 cm²) and secured using a vice. It was patterned after the design of a similar tract (intended for vowel modeling) by Takayuki Arai.⁵ The plates can be ordered such that their various apertures model the area function of any number of voiceless fricatives. In this study, the alveolar fricative [s] is investigated. The vocal tract area function for the fricative was based on an MRI study of American English fricatives (Narayanan et al. 1995). The area function for model [s], using data found in this study, is given in Figure 2.4.

Together, the drilled plates constitute a model of the oral cavity with apertures representing oral cavity constrictions during the production of the American English fricative [s]. During the production of a fricative, three variables are considered of greatest importance:

1. Dimensions of the cavity anterior to the supraglottal constriction;

⁵I express my appreciation to Professor Arai for his generous donation of this earlier model to the Berkeley Phonology Laboratory.

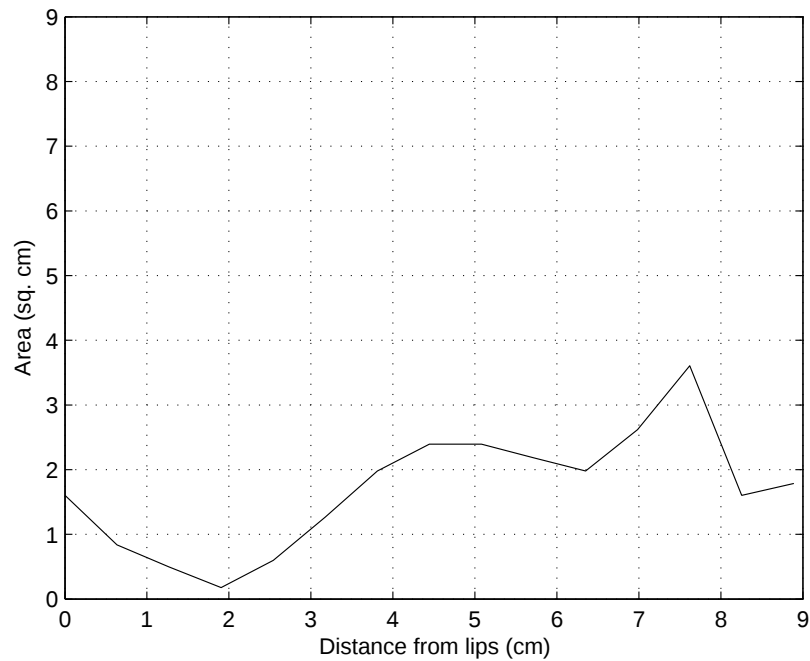


Figure 2.4: Area function of an American English alveolar fricative [s] as used in the design of the mechanical fricative model. Measurements are based on Narayanan et al. (1995).

2. Dimensions of the narrowest supraglottal constriction.
3. Presence of an obstacle or spoiler, either an edge obstacle or a wall obstacle (Shadle 1997).

Each of these components could be represented and varied in the mechanical model. The dimensions of the constrictions, as previously mentioned, were modeled by the variably-sized holes drilled in the acrylic plates. A thin acrylic plate (0.16 cm thick) mounted between the front plates of the pseudo-oral cavity (and directly in the path of the flow) served as a spoiler, i.e. a model of the incisors. As with the variable constriction sizes, the placement of the obstacle depended on the articulatory specification set forth in Narayanan et al. (1995).

In the model, the region posterior to the velum, the oropharynx, is a functional abstraction of the oropharynx, the design of which is not based on physiological measurement. The oropharynx is modeled as a sealed container that opens to the pseudo-oral cavity, an air supply, a port for the digital manometer to measure pressure, and the velopharyngeal orifice. Through one of the four holes drilled in the pseudo-oropharynx, a tube of length

60.96 cm and interior diameter 2.54 cm could be plugged with aluminum stoppers with various internal diameters. Stoppers of different diameters could be plugged into the tube to model different velopharyngeal orifice sizes, thus shunting air from the pseudo-oral cavity in a systematic manner. In practice, nine velopharyngeal orifices of different surface area were used in the experiment: 0, 0.005, 0.020, 0.045, 0.079, 0.178, 0.317, 0.495, and 0.713 cm². The tube itself was of course much longer than a typical nasal passage. This was done so that the air exiting the tube would have less influence on the sound recorded at the opening of the pseudo-oral chamber. A photograph of the model is provided in Figure 4.1 at the end of the manuscript.

Air was discharged into the pseudo-oropharynx at a constant rate from a pressurized source through a tube of length 30 cm and an interior diameter of 0.5 cm. The level of discharge was determined by trial-and-error, sampling and calibrating the pressure behind the constriction until it reached a canonical level for fricatives (8–10 cm H₂O).

2.7.2 Model data

While air was discharged into the model, pressure and audio were continuously sampled. Recordings were made at pseudo-velopharyngeal openings (VPO) ranging from 0 cm² to 0.72 cm²). During the recording, the aperture was periodically closed and re-opened. The records indicate that during the open phase, pressure dropped and the acoustic signal was attenuated accordingly (see Figure 2.5).

All data signals were digitized at 20 kHz using a Dell Optiplex GX270 computer, a multifunction data acquisition board (Model PCI-6013, National Instruments Corp., Austin, TX) with a shielded connector block (Model BNC-2110, National Instruments Corp.), and Matlab 7.0.4 software running on a Windows XP platform in the Valley Life Sciences Building, University of California at Berkeley.

Model audio Model audio was recorded using a cardioid dynamic microphone (frequency range 30 to 16,000 Hz) (Model D-190E, AKG Acoustics, Nashville, TN) and a dual microphone pre-amplifier (Model SX202, Symetrix, Inc., Mount Lake Terrace, WA). The microphone was positioned approximately 5 cm from the pseudo-oral exit of the model.

Pressure The model was connected to a pressure transducer (Model PTW-1, Glottal Enterprises, Syracuse, NY) using a tube of length 10 cm and interior diameter of 0.5 cm.

The output from the transducer was low-pass filtered (4-pole, Butterworth) at 75 Hz using an analog filter (Model 3364, Krohn-Hite Corp., Brockton, MA).

Pressure calibration A digital manometer (Model DM-1, Infiltec, Inc., Waynesboro, VA) was used to calibrate the pressure signals. Using a syringe, the electrical response of the transducer was recorded at approximately -1, 0, and 1 cm H₂O, and then related to the readings from the digital manometer using least-squares linear regression.

2.8 Acoustic analysis

2.8.1 Segmentation

For spoken fricatives, signals were manually segmented from the last glottal pulse of the vowel preceding the fricative to the first glottal pulse of the vowel following the fricative. Spectrograms were used to help determine the position of the glottal pulses.

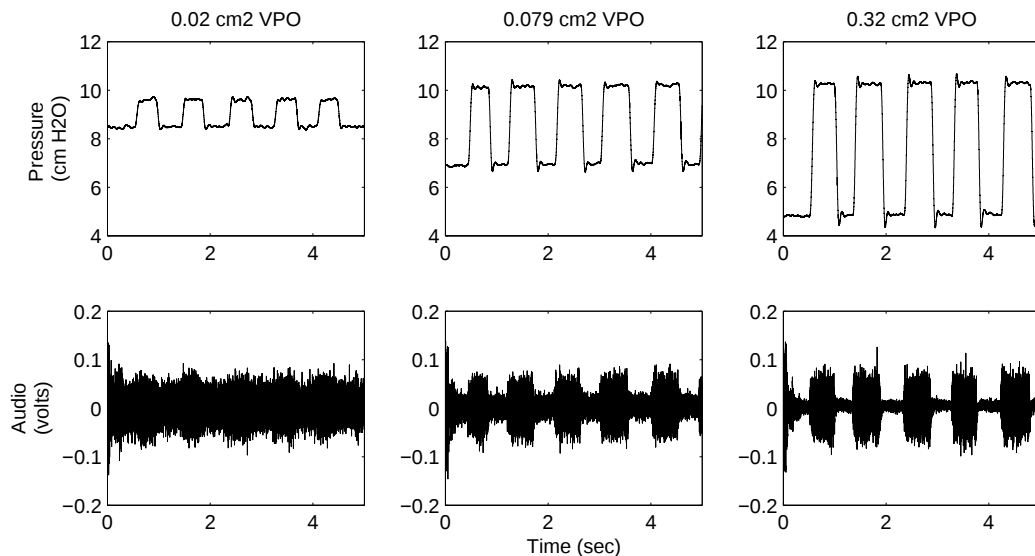


Figure 2.5: Pressure and audio recordings of the analog fricative [s] at three different pseudo-velopharyngeal openings (VPO), 0.02 cm², 0.079 cm², and 0.32 cm². Peaks in pressure represent periods when the aperture was closed, clearly accompanied by an increase in audio amplitude. These increases are predictably greater for larger VPOs.

For the mechanical fricatives, abrupt changes in the pressure signal were used as landmarks to manually segment the open (nasalized) phases of the signals. Measurements

were taken from 5 ms after opening to 5 ms before closing. Figure 2.5 illustrates a recording of the model fricative [s] at various pseudo-velopharyngeal openings (VPO). The troughs in the pressure signal (corresponding to higher amplitudes in the acoustic signal) are indicative of the open phase at various VPOs.

2.8.2 Normalization

Acoustic signals were not time-normalized.

2.8.3 Zero-crossing rate

According to (Rabiner and Schafer 1978: 127), “[A] zero-crossing is said to occur if successive samples [in a discrete-time signal] have different algebraic signs.” One simple way of measuring the frequency content of a signal is to measure the rate at which zero crossings occur. Moreover, “there is a strong correlation between zero-crossing rate and energy distribution with frequency” (Rabiner and Schafer 1978: 128). The authors further generalize that a high zero-crossing rate characterizes an unvoiced speech signal and a low zero-crossing rate characterizes a voiced one. This is due to the differing source characteristics of voiced and voiceless sounds, e.g. the reduction in airflow during the voiced sound. Nasal venting may have an analogous effect.

Zero-crossings (including crossings with both positive and negative slopes) were counted for non-normalized signals using a Matlab script (Brueckner 2002). Total zero-crossings were then divided by the duration of each fricative to determine zero-crossings per second (ZC/s), the zero-crossing rate, or ZCR.

2.8.4 Power spectra

Signals were divided into 200-point (10 ms) frames, one right-aligned, one left-aligned, one centered, and six spaced equally between the center of the edge-aligned frames and the center of the signal (three on each side of the center). Figure 2.6 illustrates the spacing of the nine frames during the fricative in the sequence [ĩĩ] (Brazilian Portuguese). In this example, there happens to be no overlap between the frames. In shorter signals, frames did in fact overlap.

Each frame was then mathematically transformed using a 200-point (10 ms) Hamming window to reduce edge-effects, as illustrated in Figure 2.8 for the center-frame data

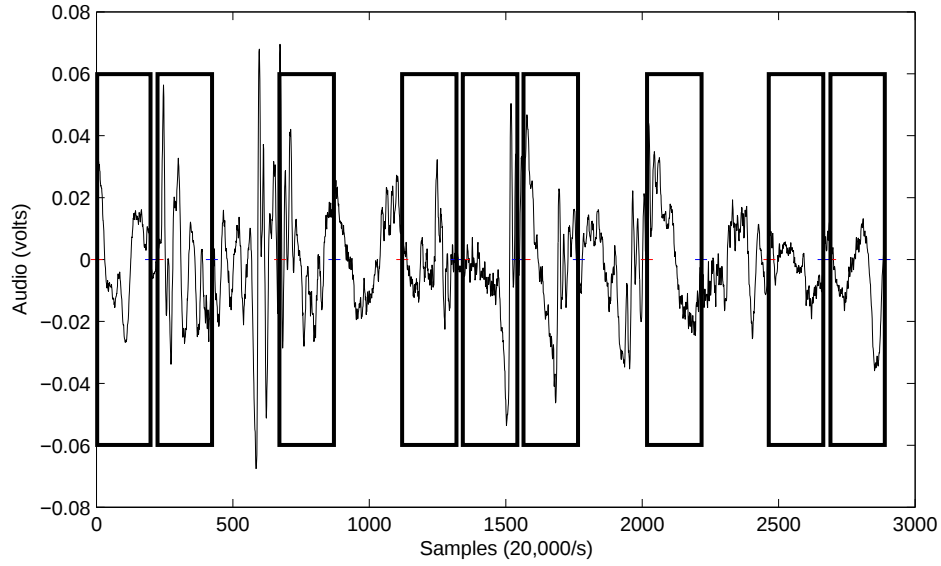


Figure 2.6: Spacing of nine 200-point (10 ms) frames applied to the audio signal $[x]$: one centered, one left-aligned, one right-aligned, and six spaced equally between the center of the edge-aligned frames and the center of the signal (three on each side of the center frame).

of $[x]$.

After application of the Hamming function to each frame, a 1024-point discrete Fourier transform (DFT) was then applied to each window. Since the windows were 200 samples long and DFTs were 1024 samples long, the windows were padded with trailing zeros to reach length 1024. The discrete Fourier transform (DFT) of the Hamming-window center-frame data appears in Figure 2.9.

Spectral averaging techniques

The methodology presented here follows closely the averaging techniques set forth in Jesus and Shadle (2002: 444–445). The authors present two techniques: time-averaging and ensemble-averaging, both of which will be used in the present analysis.

Time-averaging The time-averaged power spectrum for each fricative is given by

$$P_T(f) = \frac{1}{W} \sum_{i=1}^n |X_i(f)|^2 \quad (2.2)$$

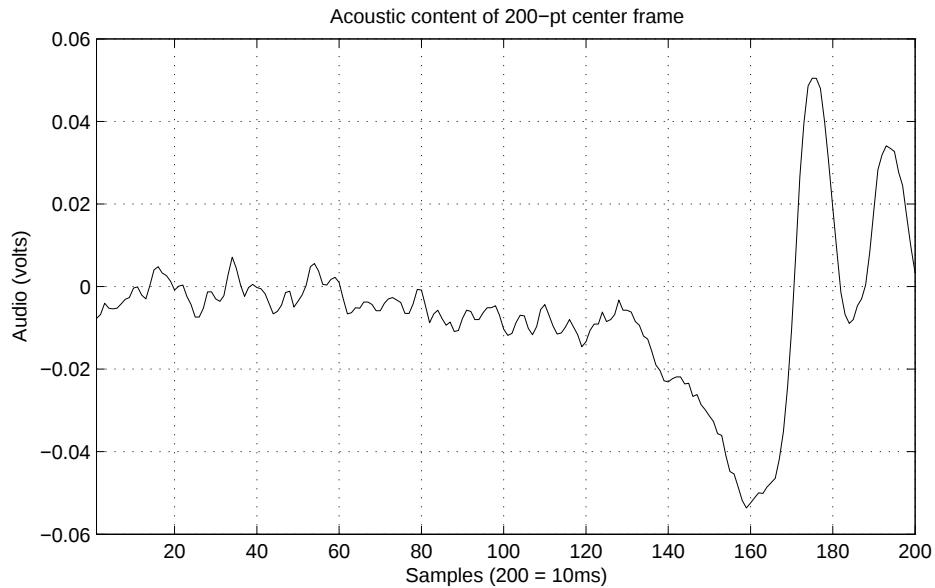


Figure 2.7: Acoustic data from the 200-point (10 ms) center frame of a velar fricative [x].

where X_i is the DFT of a portion of the fricative signal, x_i , corresponding to the i -th windowed segment of each fricative. $P_T(f)$ therefore represents the power spectrum of a given fricative, averaged across W windows ($W = 9$ for both the mechanical and spoken fricatives in this study) overlaid on the fricative. Figure 2.10 is an example of a time-averaged spectrum for a velar fricative [x].

Ensemble-averaging The ensemble-averaged power spectrum for each fricative is given by

$$P_E(f) = \frac{1}{N} \sum_{i=1}^N |X_k(f)|^2 \quad (2.3)$$

where X_k is the DFT of a portion of the fricative signal, x_k , corresponding to the windowed segment of the k -th token. $P_E(f)$ therefore represents the power spectrum of a given window, averaged across N tokens of that fricative. Here, data were usually gathered for 9–21 windows, whereas Jesus and Shadle (2002) were interested only in the acoustic properties of the beginning, middle, and end of the fricatives.

Ensemble-averaging is a useful technique for identifying the time-varying properties of fricatives and so is closely linked with coarticulation. During the production of the

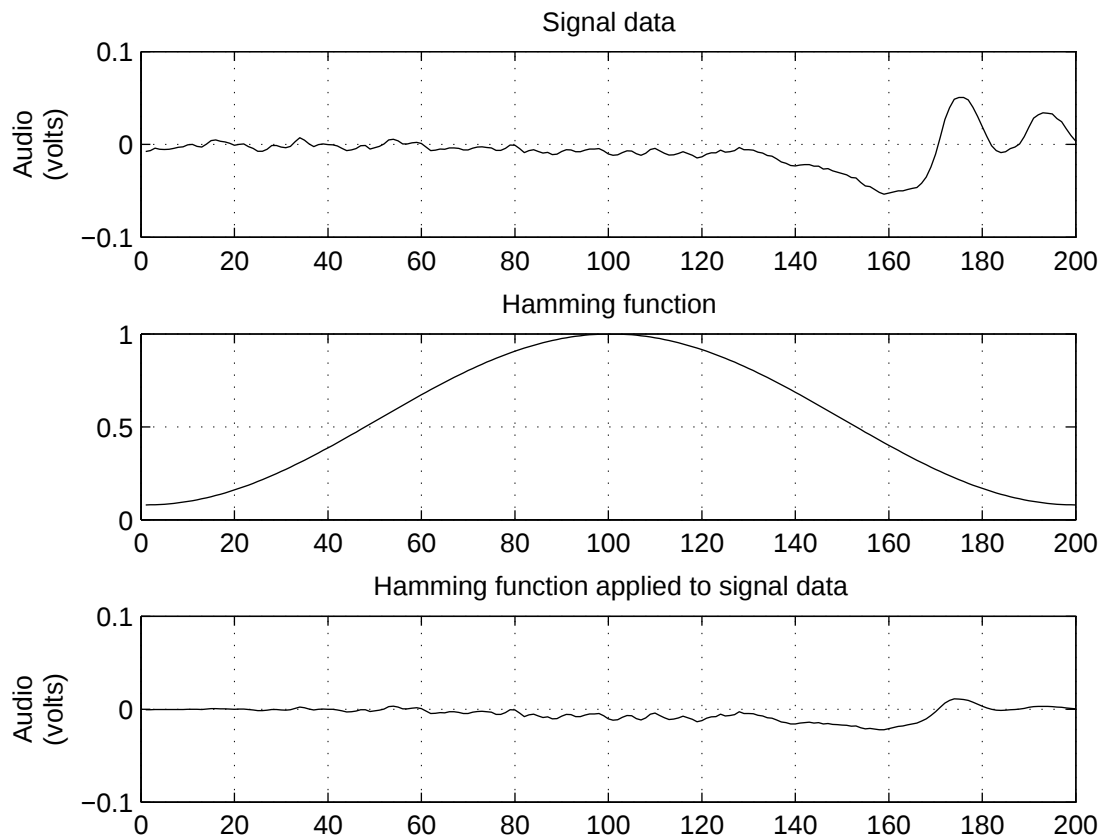


Figure 2.8: Original signal data, 200-point Hamming function, and Hamming function applied to the original signal data of the 200-point (10 ms) frame at the center of [x]. As demonstrated, the Hamming function gradually reduces the amplitude of the signal towards the edges of the window.

mechanical fricatives (Section 2.2.2) there is no *a priori* reason to believe that significant time-variation will occur, so it is not necessary to use ensemble-averaging. However, application of this technique to the mechanical data will serve as a useful point of reference when the time-averaging technique is applied to the spoken data. The degree of variation from $|X_k(f)| \cdots |X_z(f)|$ (e.g. window 1 to 9) for any given acoustic parameter should be relatively small for the time-invariant mechanical fricatives vis-à-vis the degree of variation for the time-variant (coarticulated) spoken fricatives. Accordingly, a brief look at ensemble-average results for mechanical fricatives is presented in Section 3.3.2.

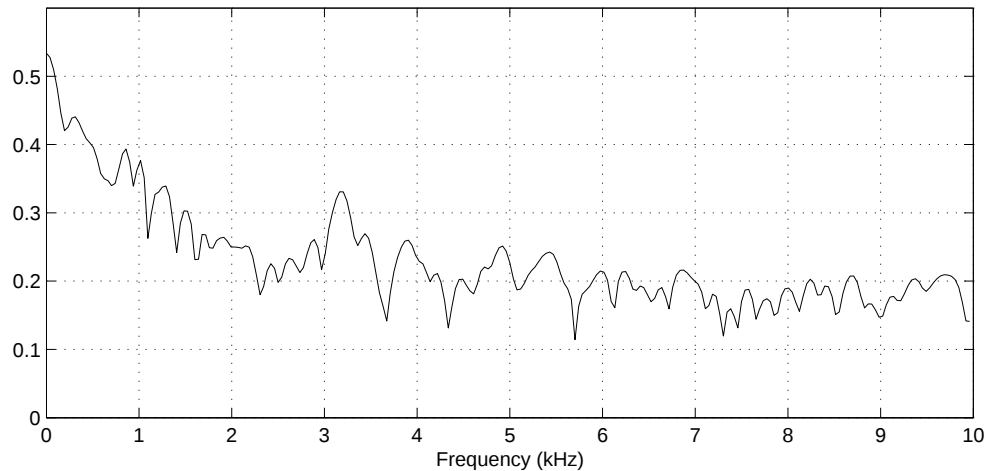


Figure 2.9: Frequency content in the central 10 ms of the fricative in [ĩĩ], pre-processed using a Hamming window, calculated using a 1024-point discrete Fourier transform.

2.8.5 Parameterization of fricative spectra

Several measures were used to extract information about the frequency content of the various fricatives and the various regions of each fricative. Following Jesus and Shadle (2002: 445–448), three important parameters were defined for each spectrum: F , $\bar{F}$, and f . The maximum spectral amplitude, F , of each signal was first established. The position of F is crucial to spectral tilt measures (High frequency and low frequency spectral slope and dynamic amplitude, as reviewed in this section) because it constitutes an endpoint for the linear regression lines use to calculate each. In practice, F was defined as the frequency with the maximum spectral amplitude occurring between 0.5 and 20 kHz. The lower bound of 0.5 kHz was set to exclude the fundamental frequency and its first few harmonics in voiced fricatives as well as room noise recorded during voiced and voiceless fricatives.

The expectation, borne out in Jesus and Shadle (2002), is that F corresponds to the frequency of the first front cavity resonance. Accordingly, F changed position based on place of articulation and vowel context. The values of F could range widely (up to 3.6 kHz for relatively flat labiodental spectra) (Jesus and Shadle 2002: 447). Because these flat-spectra variations are not of particular interest, the parameter $\bar{F}$ was computed as the average (rounded to nearest kHz) of the values of F for all tokens for each place of articulation for all speakers. Thus, there was a single $\bar{F}$ value for the fricatives of each language (e.g. Hindi [s], Portuguese [f], French [ʃ], etc.). By definition, $\bar{F}$ ignores spectral

changes based on vowel context. Analyses based on $\bar{F}$ only are presented in Chapter 3.

The third parameter f is defined as the frequency of the minimum spectral amplitude occurring between 0 and 2 kHz. The parameter f is used in the calculation of dynamic amplitude or DynAmp (Section 2.8.5).

Figure 2.10 illustrates some of these parameters for a time-averaged velar fricative [x], along with measurements to be discussed in sections below.

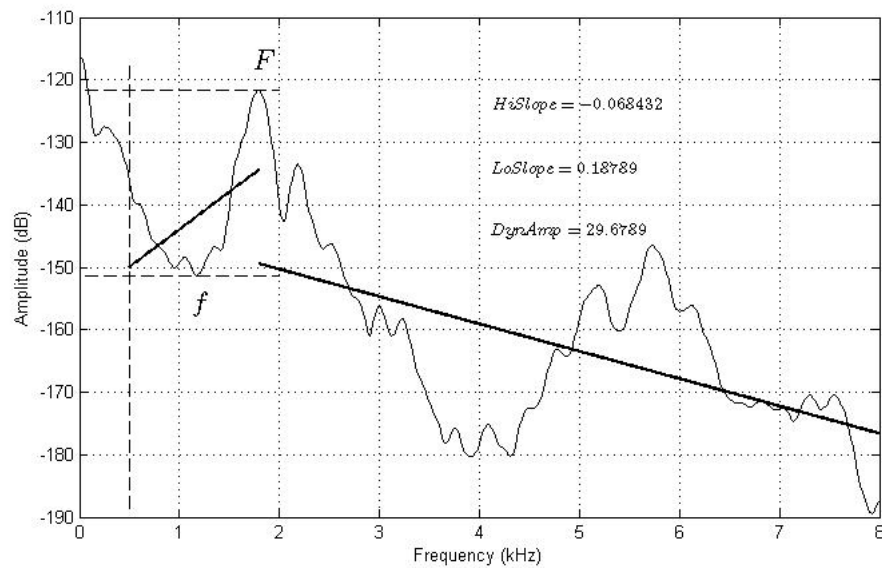


Figure 2.10: Parameterization and acoustic measurements for the time-averaged power spectrum of a velar fricative [x] (note that test tokens were sampled at 20 kHz). F is the first spectral peak occurring between 0.5 kHz and the highest frequency in the DFT (here, the sampling rate is 16 kHz). f is the minimum value between 0 and 2 KHz. Thus, $F - f$ = the dynamic amplitude or DynAmp (Section 2.8.5). HiSlope (Section 2.8.5) is the slope of the bold line on the right side of the diagram and LoSlope (Section 2.8.5) is the slope of the bold line on the left.

High frequency spectral slope (HiSlope)

This is the slope of the least-squares linear regression line fitting all points between the spectral amplitude at $\bar{F}$ and the spectral amplitude at 20 kHz. For a given fricative, high slope spectral frequency “should increase, i.e., become less negative, as flow velocity through the constriction increases” (Jesus and Shadle 2002: 448).

Low frequency spectral slope (LoSlope)

This is the slope of a least-squares regression line fitting all points between the spectral amplitude at 0.5 kHz and the spectral amplitude at $\bar{F}$. For a given fricative, low frequency spectral slope should vary directly with source strength because greater source strength ought to maximize the amplitude at F (Jesus and Shadle 2002: 448).

Slope reference

For reference, a diagram is provided that reviews the nature of slope increase and decrease (Figure 2.11).

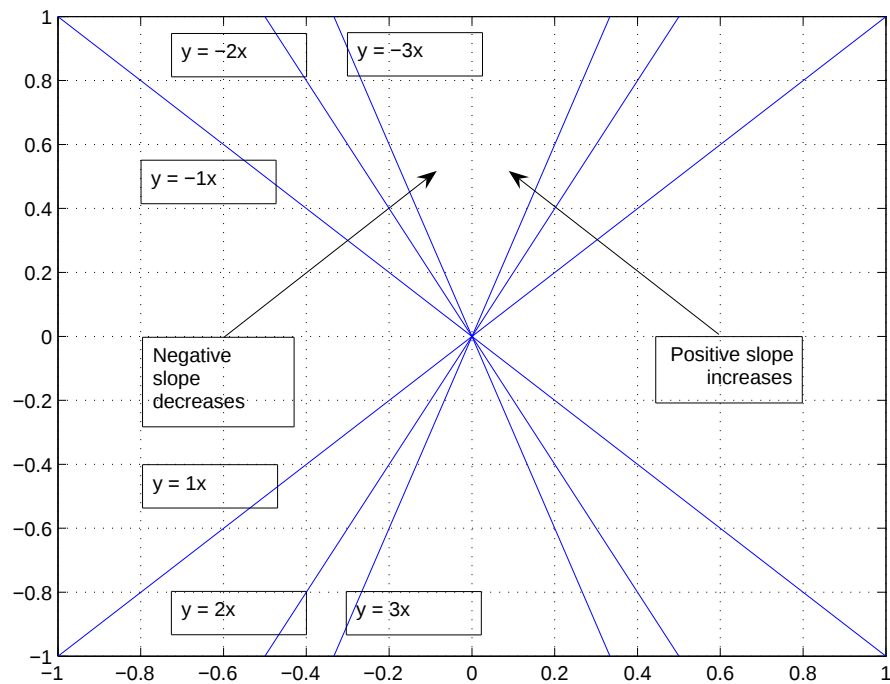


Figure 2.11: A slope diagram. The measure HiSlope (typically negative slope) is expected to *increase* (i.e. become less negative) for nasalized fricatives. LoSlope (typically positive) is expected to *decrease*.

HiSlope is expected to be negative since spectral energy should be falling. Thus, when negative slope is decreasing, it means a slope that falls more steeply, e.g. a more precipitous decline in spectral energy. When negative slope is increasing, the spectrum is

flatter, or there is less energy in the high frequency range. The inverse is true of positive slope. LoSlope is expected to be positive because it measures the rise to the first spectral peak (from 0.5 kHz). A *decrease* in positive slope indicates a steeper rise to the spectral peak whereas a *decrease* indicates flatness in the low frequency part of the spectrum.

In terms of the research hypothesis, it is expected that both HiSlope (positive slope) and LoSlope (negative slope) will be smaller for nasalized fricatives than for oral fricatives.

Dynamic amplitude (DynAmp)

This value represents the difference between the maximum amplitude of the spectrum occurring between 0.5 kHz and 10 kHz and the minimum amplitude occurring between 0 and 2 kHz. According to Jesus and Shadle (2002: 448), this parameter “should be maximized for a localized source, and for higher relative noise source strength, as in sibilants and unvoiced fricatives.” Analogously, it is expected that the reduction in noise source strength caused by velopharyngeal insufficiency and/or nasalization during a given fricative should reduce dynamic amplitude. This measure is not expected to be very large for fricatives with relatively flat spectra, such as labiodentals.

High wide-band frequency energy (HiBand)

This is a static measure of the average spectral amplitude found between 3.5 and 6 kHz.

Spectral peak bandwidth

While all the preceding measures were based on FFT-analysis, the measure of spectral peak bandwidth is based on LPC-analysis. Fourteen coefficients were used to detect peaks in the fricative spectrum (using the frame-alignment techniques discussed in Section 2.8.4). The width of the first spectral peak occurring above 150 Hz was measured at a depth of 4 dB and is reported in Hz.

2.9 Flow analysis (spoken fricatives)

2.9.1 Segmentation

Aerodynamic signals were segmented in tandem with the acoustic signals discussed above. Thus, the start- and end-points of the oral and nasal flow signals, as well as the pressure signals for the analog fricatives, corresponded exactly to those of the parallel acoustic signals.

2.9.2 Normalization

Each signal was segmented into 100 equally-spaced intervals and an average value was computed for each. Thus, each normalized signal was comprised of exactly 100 samples. Though it reduced the data resolution for the average signal, normalization was a necessary step for undertaking the polynomial fitting and numerical integration of the signals (see Sections 2.9.3 and 2.9.4).

2.9.3 Polynomial fitting

Coefficients

A third-degree polynomial $f(x)$ that fits the time-normalized aerodynamic signals in a least-squares sense was calculated for each aerodynamic signal, using Matlab 7.0.4. The algorithm forms the Vandermonde matrix,⁶ V , whose elements are powers of x , where

$$v_{i,j} = x_j^n - j \quad (2.4)$$

The algorithm then solves the least squares problem $V_p \cong y$ for each Vandermonde matrix.

Cubic polynomials were selected because their characteristic shape models the oral flow pattern for fricatives which tend to consist, maximally, of a peak, a valley, and a peak. Similarly, for nasal flow (in the $\tilde{V}C\tilde{V}$ context) there will be a peak, a valley and a peak, where the peaks correspond to the nasal vowels and the valley corresponds to the fricative. The top frame of Figure 2.12 illustrates nasal flow during the fricative in the sequence $[\tilde{a}f\tilde{a}]$ (Hindi). A cubic polynomial has been fitted to the aerodynamic data. The four coefficients

⁶Vandermonde matrices are a useful tool in polynomial interpolation precisely because solving the system of linear equations $Vu = y$ for u (where V is the $n \times n$ Vandermonde matrix) is the same as finding the coefficients u_j of the polynomial $P(x) = \sum_{j=0}^{n-1} u_j x^j$ of degree $\leq n-1$ which has values y_i at α_i .

of the equation are given at the top of the Figure 2.12. The same is true for Figure 2.13, only that the same procedure has been invoked to analyze the oral flow from the same token.

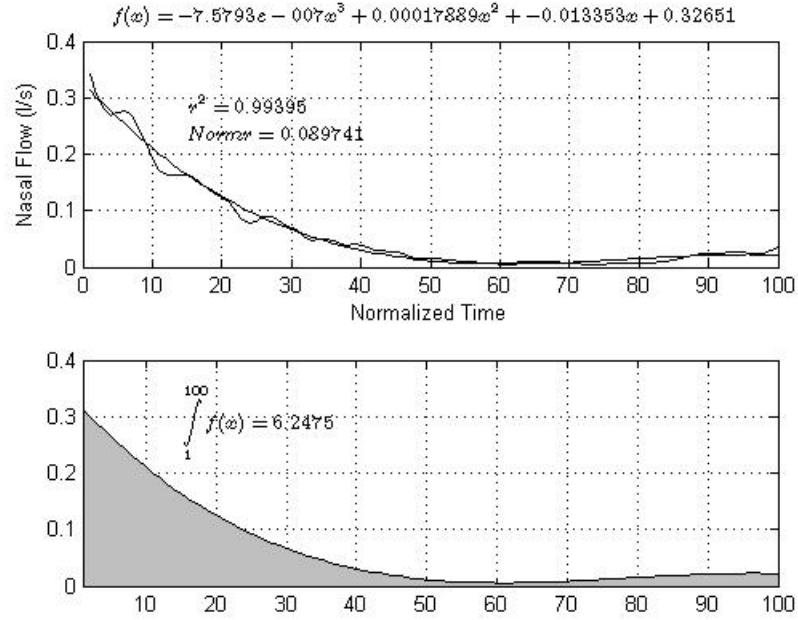


Figure 2.12: Top frame: Nasal flow during the fricative in [ãfã] (Hindi) and a third-degree polynomial fitted to the nasal flow. Bottom frame: The shaded portion represents the numerical integral of the flow $\int_1^{100} f(x) = 6.2475$.

Correlation

The correlation coefficient of the normalized signal data and the cubic polynomial were computed. The normalized signal data and polynomial in Figure 2.12 have a correlation of 0.99395. In other words the polynomial function accounts for approximately 99% of the normalized signal data.

Norm of residuals

A norm of residuals was calculated for the cubic polynomial fit to each normalized signal. The norm of residuals of the (time-normalized) recorded data versus the polynomial

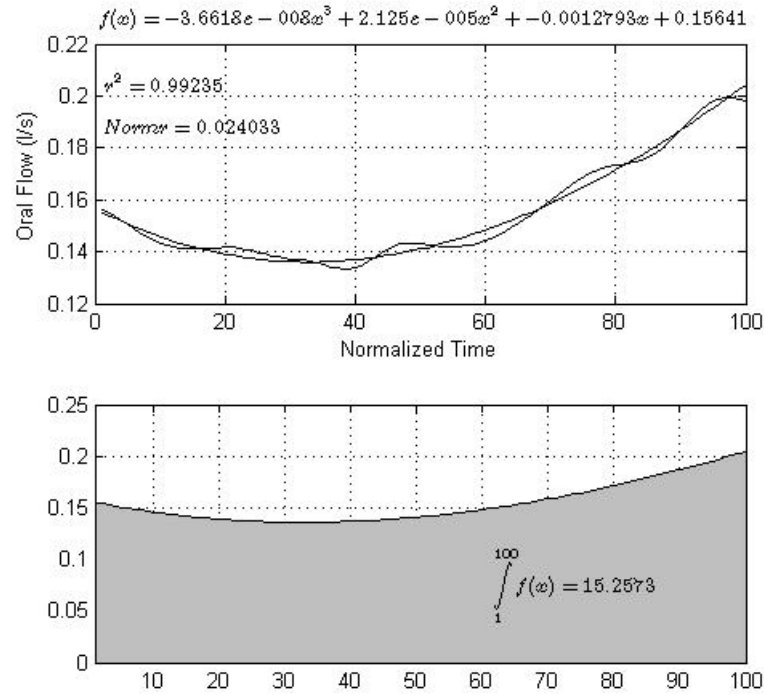


Figure 2.13: Top frame: Oral flow during the fricative in [āfā] (Hindi) and a third-degree polynomial fitted to the oral flow. Bottom frame: The shaded portion represents the numerical integral of the flow $\int_1^{100} f(x) = 15.2573$.

fit is defined as:

$$\sqrt{\sum_{i=1}^n (Y_i - Yp_i)^2} \quad (2.5)$$

where Y is the (time-normalized) aerodynamic recording Yp is the best-fitting cubic polynomial.

Figure 2.14 illustrates the norm of residuals between the cubic polynomial and the data in Figure 2.12 above. The total norm of residuals, 0.089741, is the sum of the residuals at each data point across the normalized x-axis.

Statistical evaluation of polynomial fit

Tokens in which either the oral or nasal polynomial fit had a norm of residuals greater than three standard deviations from the mean or a correlation coefficient r greater than three standard deviations below the mean were excluded from further statistical anal-

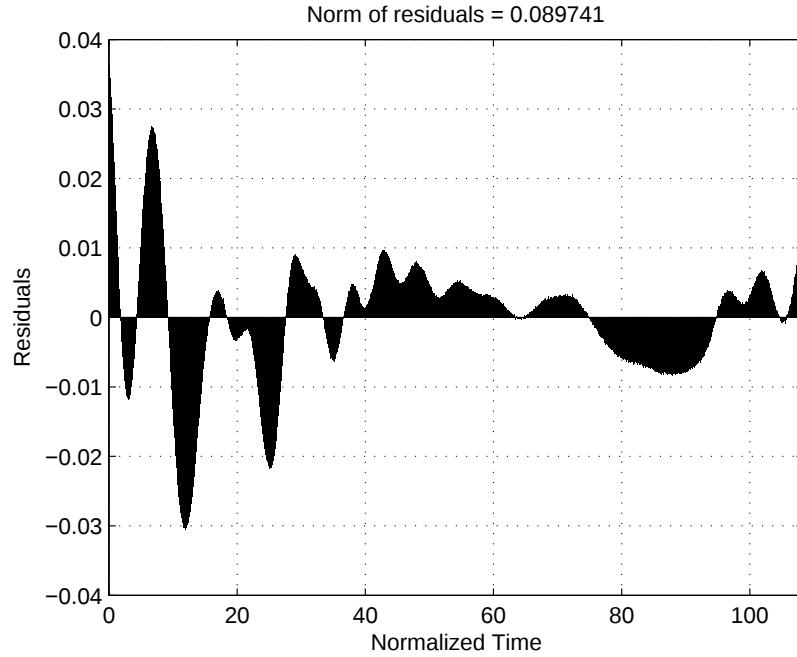


Figure 2.14: Residuals for the cubic polynomial fitted to nasal flow during the fricative in [ãxã] (Hindi).

ysis. Such tokens were considered outliers. For such tokens, it was judged that a cubic polynomial could not reasonably approximate the normalized airflow geometry of the fricative.

2.9.4 Numerical integration

Using Matlab 7.0.4, the polynomial coefficients for each aerodynamic signal were passed to anonymous functions. These functions were then fed into a numerical integration algorithm that tries to approximate the integral of a function from a to b (the start- and end-points determined by acoustic segmentation to within an error of $1e-6$ using recursive adaptive Simpson quadrature (Gander and Gautschi 2000). If we compute the value of some integral

$$\int_a^b f(x)dx = I(f) \quad (2.6)$$

to within a given error tolerance, we generally use a standard quadrature formula, such as Trapezoidal rule. Under this regime, the ‘worst behavior’ of the function determines the

dimensions of the grid. To approximate that portion of the integral where the function varies rapidly (or ‘behaves badly’) we overlay a sufficiently fine grid to account for the variation. When the variation decreases, however, a coarser grid may be used. An adaptive procedure like Simpson quadrature automatically chooses a nonuniform grid in order to approximate the integral of the function within a specified error tolerance and with the greatest degree of efficiency.

Integrals were approximated for integrands corresponding to both the oral and nasal flow signals of each token. The resulting values, approximations of the areas beneath the curves of the normalized airflow signals, were taken to be holistic estimates of nasal flow and oral flow during the production of each fricative token. The bottom frames of Figures 2.12 and 2.13 illustrate the calculated areas beneath two integrands. In the case of the two figures, cubic polynomials have been fitted to time-normalized nasal and oral flow during the production of the fricative in the sequence [ãfã] (Hindi). The numeric approximations of the integrals, according to adaptive Simpson quadrature, are given in the figures themselves.

2.9.5 Maximal flow rate and flow rate at temporal center

Maximum values (in l/s) were tabulated for the oral and nasal signals. The measure of flow at the temporal center of the aerodynamic signal was also tabulated.

2.10 Pressure analysis (mechanical fricatives)

After the pressure signals had been segmented as described in Section 2.9.1, pressure (in cm H₂O) was averaged across the excised signal.

2.11 Statistical Methods

2.11.1 Review of variables

Continuous variables

The continuous acoustic variables are reviewed in Table 2.1.

Categorical variables

There are four categorical variables, reviewed in Table 2.3.

Table 2.1: Continuous acoustic variables

Continuous variable: Acoustic	Described in	Applies to data type(s)
Zero-crossing rate (zc/s)	2.8.3	mechanical & spoken
High frequency slope (dB/kHz)	2.8.5	mechanical & spoken
Low frequency slope (dB/kHz)	2.8.5	mechanical & spoken
Dynamic amplitude (dB)	2.8.5	mechanical & spoken
High wide-band frequency energy (kHz)	2.8.5	mechanical & spoken

Table 2.2: Continuous aerodynamic variables

Continuous variables: Aerodynamic	Described in	Applies to data type(s)
Flow equation integrals	2.9.4	spoken only
Flow maxima (l/s)	2.9.5	spoken only
Flow temporal center (l/s)	2.9.5	spoken only
Pressure (cm H ₂ O)	2.7.2	mechanical only
Pseudo-velopharyngeal aperture (cm ²)	2.7.1	mechanical only

2.11.2 Null hypotheses

Spoken fricatives

The null hypotheses for spoken fricatives are as follows:

1. The means of aerodynamic measures for spoken fricatives (see Table 2.1) will not differ significantly based on nasal context, i.e. whether the fricatives are uttered in VCV or $\tilde{V}C\tilde{V}$ syllables.
2. The means of acoustic measures (see Table 2.2) will not differ significantly based on nasal context.

In other words, the experiment will attempt to show that fricatives differ in their spectral and aerodynamic properties when they are under the effects of coarticulatory nasalization.

Table 2.3: Categorical variables

Categorical variables	Described in	Applies to data type(s)
Nasal control group	2.5	spoken only
Language	2.3	spoken only
Speaker	2.4	spoken only
Place of articulation	2.7.1	mechanical & natural

Mechanical fricatives

The null hypotheses for mechanical fricatives are the following:

1. Pressure (cm H₂O) under the effects of differing pseudo-velopharyngeal apertures share a common mean;⁷
2. Acoustic measures (see Table 2.1) under the effects of differing pseudo-velopharyngeal apertures (in cm²) share a common mean.

That is to say, the results of the experiment will show whether or not there is a significant relationship between the size of a model velo-pharyngeal vent and various acoustic measures that seem important to the acoustics and perception of fricative sounds.

2.11.3 Linear statistical models

Variables that could reasonably be assumed to have a normal distribution were incorporated in linear models. The normal distribution characteristics of each continuous variable were assessed using the Lilliefors test, described below. In cases where variables failed either test of normality, the data were either transformed as described below or, failing acceptable results, incorporated in non-linear models.

Normality

Lilliefors test This test is similar to Kolmogorov-Smirnov but instead of comparing the distribution of the given variable to a standard normal distribution, the Lilliefors test compares the empirical distribution of the variable with a normal distribution having the same mean and variance as the variable itself (Lilliefors 1967). Indeed, Lilliefors adjusts for the fact that the parameters of the normal distribution are estimated from the given variable rather than specified in advance. The result 1 indicates that we can reject the hypothesis that the variable has a standard normal distribution. The result 0 indicates that we cannot reject that hypothesis. In the present study, the null hypothesis is rejected if the test is significant at the 0.05 level.

⁷While the research hypothesis, i.e. that they will *not* share a common mean, is an accepted and indeed fundamental principle of aerodynamics, nonetheless it seems prudent to demonstrate the effect for present purposes.

Data transformations When variables failed Lilliefors, they were mathematically transformed according to guidelines set forth in (Hoaglin and Hoaglin 1981). Right-skewed data (clustered at lower values) were transformed using lower-power transformations (e.g. square root, cube root, logarithmic transformations, etc.). Left-skewed data (clustered at higher values) were transformed using higher-power transformations (e.g. cube, square, etc.). Lilliefors was used again to assess the normality of the transformed data.

One-way analysis of variance

The null hypotheses were assessed using one-way analysis of variance. Each acoustic measure is regarded individually. Tukey's honestly significant difference criterion (optimal for one-way ANOVA with equal sample sizes) is used to determine which differences are significant at the 0.05, 0.01, and 0.001 levels.

2.11.4 Non-linear models: Kruskal-Wallis

This test is a nonparametric version of one-way analysis of variance, the assumption being that the measurements come from a continuous distribution that is not necessarily normal. The Kruskal-Wallis test is based on an analysis of variance using the ranks of the data values rather than the actual data values (Hollander and Wolfe 1973, Kruskal and Wallis 1952). Tukey's honestly significant difference criterion (optimal for one-way ANOVA with equal sample sizes) is used to determine which differences are significant at the 0.05, 0.01, and 0.001 levels.

Chapter 3

Results

3.1 Overview of the results

Aerodynamic measures strongly suggest that fricatives can undergo coarticulatory nasalization. Nasal flow measures are significantly greater during fricatives in nasal ($\tilde{VC}\tilde{V}$) syllables. Oral flow means are often significantly lower in the same context. Moreover, acoustic measures indicate that this nasalization has potentially debilitating ramifications on the perception of the fricatives themselves. High energy frequency was found to fall for the fricatives produced under nasal conditions. Also, the bandwidth of spectral peaks was found to increase in the nasal syllables.

3.2 Spoken fricatives

3.2.1 Aerodynamic results

One of the fundamental questions of this study, hinted at in languages like Icelandic (Pétursson 1973, Einarsson 1940) and confirmed observationally in Coatzospan Mixtec (Geffen 1999, 2001) is this: Are fricatives between nasal vowels nasalized to any significant degree? The results of tests presented here suggest that they are.

After the calibrated nasal curves were fitted with polynomials, the integrals were compared for fricatives under the nasal and oral conditions (Section 2.9.4). The integrals themselves are rather abstract objects of comparison but they are, crucially, comparable across tokens and speakers.

Data from one speaker from each language has been used for the aerodynamic analysis. Furthermore, the population of nasal and oral fricatives was slightly reduced when correlation coefficients and norms of residuals for the polynomials showed them to be poor fits to the (time-normalized data).¹ The numbers of fricatives analyzed aerodynamically are presented in Table 3.1.

Table 3.1: Raw numbers of fricatives analyzed aerodynamically. Nasal tokens appear on the left, oral tokens on the right

	Fricative			
Language	s	ʃ	f	x
Hindi	18, 18	18, 16	18, 18	0, 0
BP	15, 18	18, 18	18, 17	18, 18
French	18, 18	17, 18	18, 16	0, 0
Totals	51, 54	53, 52	54, 51	18, 18

Mean values of aerodynamic measures for the various fricatives are presented in Table 3.2.² ‘Max’ refers to the maximum flow recorded during the fricative and ‘TC’ refers to the flow value at the temporal center of the fricative (in liters/second, see Section 2.9.5). ‘Int’ refers to the numeric integral of flow calculated throughout the duration of the fricative (see Section 2.9.4).

Table 3.2: Mean values for aerodynamic measures. Values for nasalized context ($\tilde{VC}\tilde{V}$) appear at the left of the comma, oral context (VCV) at the right. Max and TC measures are in liters/second.

Language	Nas Int	Nas Max	Nas TC	Ora Int	Ora Max	Ora TC
Hindi	13.10, 0.00	0.23, -0.10	0.06, 0.013	22.61, 51.20	0.26, 0.80	0.08, 0.44
BP	2.8, 0.83	0.14, 0.10	0.03, 0.00	5.36, 10.67	0.02, 0.21	0.07, 0.10
French	5.90, 0.14	0.23, -0.20	-0.02, -0.02	19.32, 56.46	0.40, 0.75	0.14, 0.55

Tables 3.3 and 3.4 report the F -statistics and p -values resulting from a one-way ANOVA with the various aerodynamic measures as dependent variables and nasal context as independent variable. Results are given for each language individually and for all languages

¹Approximately 5% of the tokens were discarded for these reasons.

²Negative values are likely the result of measurement error, either due to the calibration or the actual performance of the transducers. Generally speaking, they may be equated with zero nasal flow. If the flow is truly negative, the only possible physiological explanation is that the volume of the nasal cavity is somehow rarefied, perhaps due to the action of the soft palate. It is not clear what may motivate such nasal flow. Since the effect is not particularly robust, it will not be investigated further at this time. As far as the present study is concerned, it is enough to note a statistically significant relative difference between the aerodynamic measures under categorically-variable conditions.

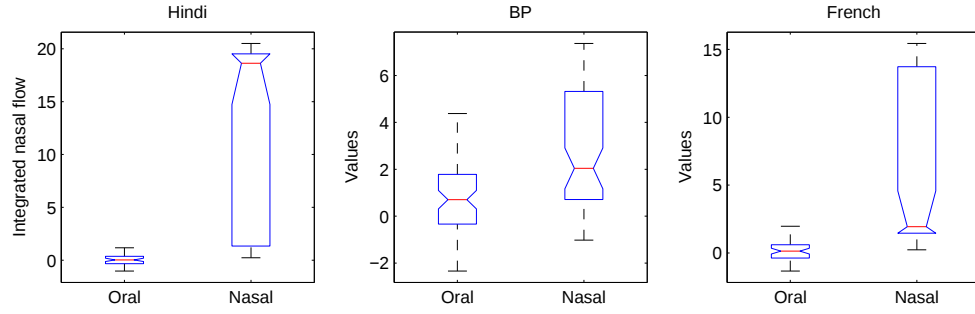


Figure 3.1: Boxplot of integrated nasal flow produced during fricatives in the nasal context $\tilde{VC}\tilde{V}$ and the oral context VCV . For Hindi $F(1, 104) = 119.05$, $p < 0.001$; for BP $F(1, 138) = 32.00$, $p < 0.001$; for French $F(1, 103) = 46.04$, $p < 0.001$.

collectively.

Table 3.3: ANOVA results for nasal aerodynamic measures by nasal context ($p < 0.05 = *$; $p < 0.01 = **$; $p < 0.001 = ***$).

Language	Nas Int	Nas Max	Nas TC
Hindi	$F(1, 104) = 119.05***$	$13.96***$	0.52
BP	$F(1, 138) = 32.00***$	0.10	0.25
French	$F(1, 103) = 46.04***$	6.16^*	0.00

Table 3.4: ANOVA results for oral aerodynamic measures by nasal context ($p < 0.05 = *$; $p < 0.01 = **$; $p < 0.001 = ***$).

Language	Ora Int	Ora Max	Ora TC
Hindi	$F(1, 104) = 23.99***$	$27.43***$	$21.04***$
BP	$F(1, 138) = 40.62***$	2.41	0.58
French	$F(1, 103) = 198.75***$	6.73^*	$83.56***$

Nasal measures

Integrated nasal flow The integrated measure of nasal flow proved significant ($p < 0.001$) in each individual language, as reported in Table 3.3. This suggests that fricatives differ from each other with respect to integrated nasal flow when they occur in nasal ($\tilde{VC}\tilde{V}$) and oral (VCV) contexts. Boxplots of these results for each individual language appear in Figure 3.1.

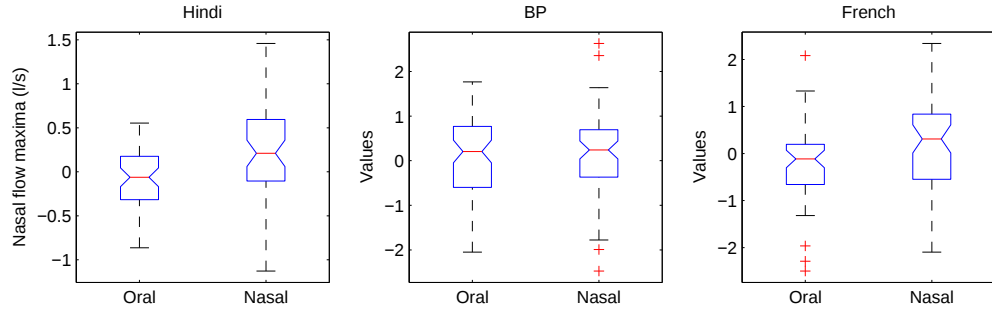


Figure 3.2: Boxplot of nasal flow maxima (l/s) produced during fricatives in the nasal context $\tilde{V}C\tilde{V}$ and the oral context VCV. For Hindi $F(1, 104) = 13.96$, $p < 0.001$; for BP $F(1, 138) = 0.10$, $p > 0.05$; for French $F(1, 103) = 6.16$, $p < 0.05$.

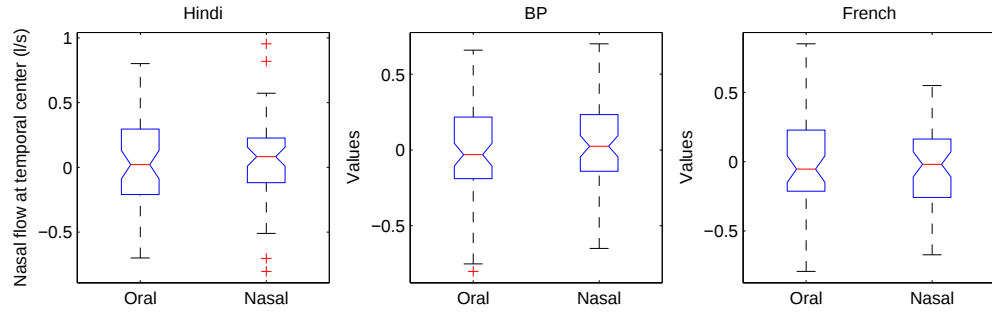


Figure 3.3: Boxplot of nasal flow (l/s) at temporal center of fricative produced in the nasal context $\tilde{V}C\tilde{V}$ and the oral context VCV. The effect is not significant for any language, $p > 0.05$.

Nasal flow maxima Maximum nasal flow, measured in liters/second, significantly differentiates fricatives occurring in nasal and oral contexts for Hindi ($p < 0.001$), and marginally for French ($p < 0.05$). The effect does not achieve significance for Brazilian Portuguese (see Table 3.3). Figure 3.2 shows the relationship between the distributions of nasal flow maxima in both contexts, for each language.

Nasal flow at temporal center The measure of nasal flow at the temporal center of the token (in liters/second), is not a significant predictor of environment for any language ($p > 0.05$). It seems unlikely that this measure could be used to reliably differentiate fricatives occurring in nasal and oral contexts. A boxplot showing the results for each language is given in Figure 3.3.

Oral measures

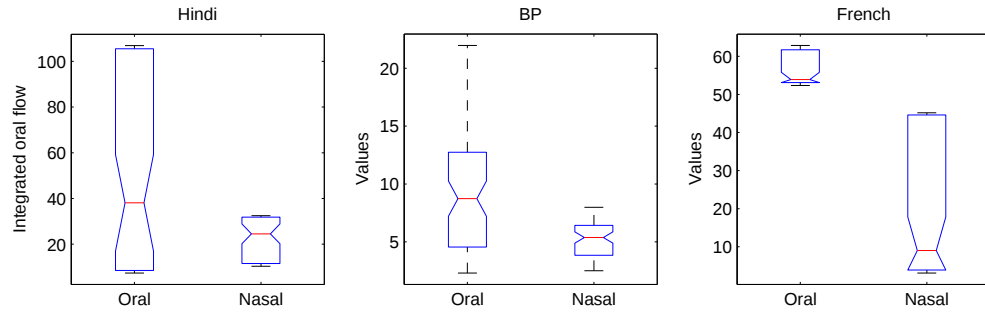


Figure 3.4: Boxplot of integrated oral flow produced during fricatives in the nasal context $\tilde{V}C\tilde{V}$ and the oral context VCV. For Hindi $F(1, 104) = 23.99$, $p < 0.001$; for BP $F(1, 138) = 40.62$, $p < 0.001$; for French $F(1, 103) = 198.75$, $p < 0.001$.

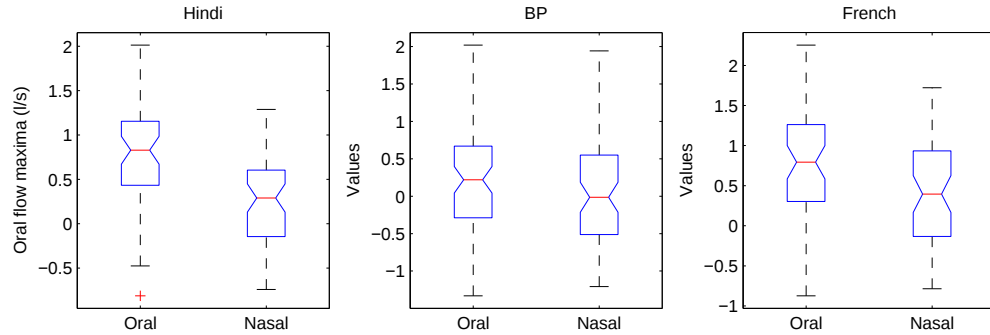


Figure 3.5: Boxplot of oral flow maxima produced during fricatives in the nasal context $\tilde{V}C\tilde{V}$ and the oral context VCV. For Hindi $F(1, 104) = 27.43$, $p < 0.001$; for BP $F(1, 138) = 2.41$, $p > 0.05$; for French $F(1, 103) = 6.73$, $p < 0.05$.

Integrated oral flow As shown in Table 3.2, integrated oral flow is consistently greater for fricatives in oral contexts (VCV) than nasal contexts ($\tilde{V}C\tilde{V}$) in each language. This effect achieves significance ($p < 0.001$) for all languages individually as demonstrated in Table 3.4. Figure 3.4 illustrates the distributions of this variable in the oral and nasal context for each language.

Oral flow maxima In some cases, oral flow maxima tend to increase for oral fricatives, vis-à-vis fricatives occurring in nasal contexts (see Table 3.2). Table 3.4 indicates that this effect is statistically significant for Hindi ($p < 0.001$) and marginally so for French ($p < 0.05$). The effect does not achieve significance for Brazilian Portuguese. The distributions for oral flow maxima in the two contexts (for each language) are given in Figure 3.5.

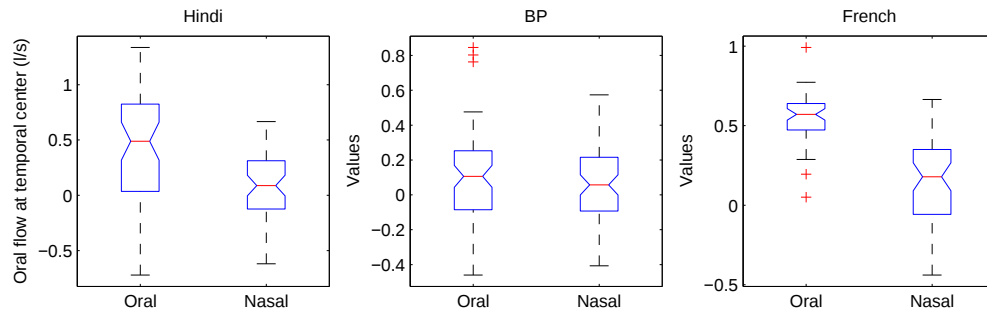


Figure 3.6: Boxplot of oral flow at temporal center of fricative produced during fricatives in the nasal context $\tilde{V}C\tilde{V}$ and the oral context VCV . For Hindi, $F(1, 104) = 21.04$, $p < 0.001$; for BP $F(1, 138) = 0.58$, $p < 0.001$; for French $F(1, 103) = 83.56$, $p < 0.001$.

Oral flow at temporal center This measure, taken at the temporal midpoint of the fricative, differs significantly ($p < 0.001$) for all languages except Brazilian Portuguese. Thus, the oral flow at this moment is typically greater for oral fricatives than it is for fricatives in nasalized contexts. In the boxplot found in Figure 3.6, the strength of the effect can be seen in each language.

Vowel context and flow measures Some readers may be interested in the effect of vowel quality on the various aerodynamic measures. Performed for each language, one-way ANOVAs showed significant results only for the measure of maximum oral flow, where the highest degree of airflow was typically found in fricatives preceded by the low back vowel. Tukey’s HSD could not differentiate between the high front and high back vowels. There was no discernible effect for V2. Boxplots of Maximum Oral Flow by vowel are presented in Figure 3.7 for each language.

These results are in line with those presented by Shosted and Willgohe (2006: 19). The authors examine the aerodynamics of voiced and voiceless stops, as well as nasals, when they occur between the three corner vowels [a i u] in Spanish. For voiced stops (which routinely spirantize in intervocalic position), they found that oral flow *minima* were greatest with a low vowel in V1 position. They attribute this difference to increased jaw opening.

3.2.2 Acoustic results

Section 3.2.1 establishes that fricatives differ significantly in terms of nasal exhalation when they are adjoined by nasal versus oral vowels. This allows us to move forward

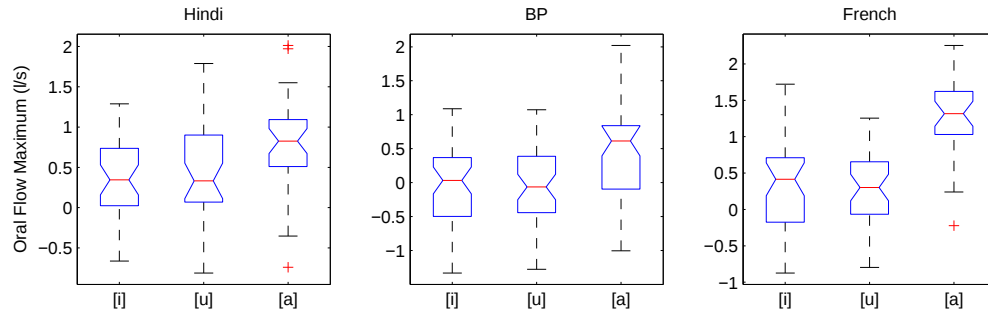


Figure 3.7: Boxplots of oral flow maxima (l/s) by vowel for each language. For Hindi $F(2, 103) = 6.45$, $p < 0.01$; for BP $F(2, 137) = 10.94$, $p < 0.001$; for French $F(2, 102) = 32.49$, $p < 0.001$. Tukey’s HSD reveals significant differences between the low vowel and the high vowels in each case.

to an acoustic analysis of the phonetically ‘nasalized’ fricatives. The central question is what makes a phonetically nasalized fricative different from a non-nasalized fricative. A secondary—though important—question is whether or not these acoustic differences are likely to be perceptible. As explained in Section 2.6.1, high quality audio was recorded in an audiometric booth when the aerodynamic masks were removed.

$\bar{F}$ by fricative For an explanation of this measure, see Section 2.8.5. According to Jesus and Shadle’s (2002) prediction, $\bar{F}$ should be lower for more posterior place of articulation. The opposite was found to be true in the present study. $\bar{F}$ of the anterior fricatives [s f] were significantly lower than those of the relatively more posterior fricatives [ʃ x]. This discrepancy may stem from the fact that fricatives were produced by speakers of different languages (none of which were European Portuguese, as in (Jesus and Shadle 2002)), where subtle articulatory differences may have affected the location of $\bar{F}$. In the present study, no predictions were made about the relation of $\bar{F}$ to nasality condition, so the discrepancy between the two studies may be overlooked for the time being. Whether or not $\bar{F}$ always behaves in the manner predicted by Jesus and Shadle (2002) with regard to place of articulation is still an open question. For present purposes, it is enough to observe that $\bar{F}$ is of significance in predicting a fricative’s place of articulation, and that posterior fricatives significantly pattern against anterior ones, though not in the anticipated direction.

$\bar{F}$ measures were significantly different ($p < 0.01$) across fricative (place of articulation) for all speakers except Hindi Speaker 2. Naturally, when the data from each speaker was pooled, the differences proved significant, as well: $F(3, 609) = 33.41$, $p < 0.01$.

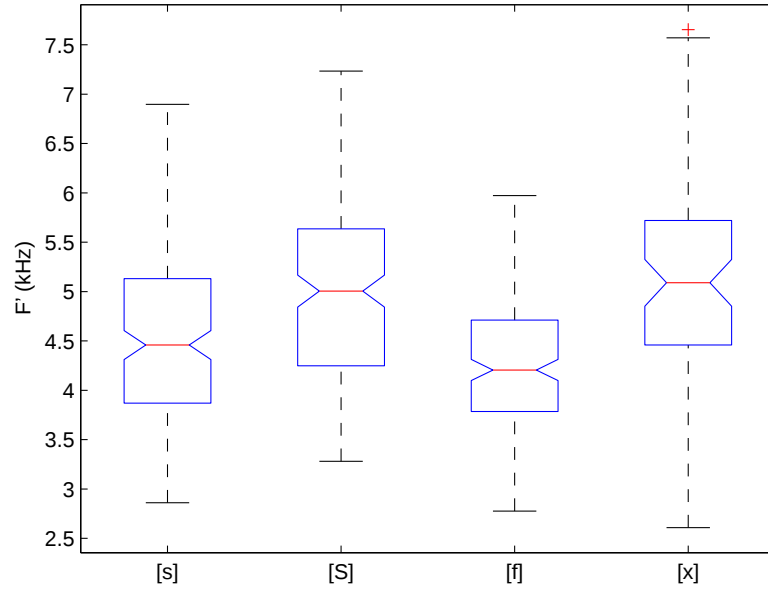


Figure 3.8: Boxplots of $\bar{F}$ values (kHz) by fricative (place of articulation). $F(3, 609) = 33.41$, $p < 0.01$. Tukey's HSD reveals the following: [s] and [f] are significantly different from each other and the rest of the fricatives; [ʃ] and [x] are significantly different from [s] and [f], but not from each other.

Boxplots of the pooled data are presented in Figure 3.8. Tukey's HSD reveals that [ʃ] and [x] are not significantly different from one another (both have a high $\bar{F}$). While [s] and [f] can be reliably differentiated from each other and from [x] and [ʃ] as well ($p < 0.05$).

Vowel quality of V1 could also be used to predict the $\bar{F}$ of the fricatives ($F(4, 608) = 2.62$, $p < 0.05$), suggesting a significant degree of coarticulation. Not surprisingly, the first frame of the fricative was most sensitive to the coarticulatory effect of V1 ($F(4, 608) = 4.46$, $p < 0.01$). Tukey's HSD suggests that the significant difference lies between [a] and the high vowel pair [i u]. $\bar{F}$ values for [ɛ ɔ] are not significantly different ($p > 0.05$) from either the low or high vowels.

The nasality of the following or preceding vowel was not a good predictor of $\bar{F}$ for any speakers.

Zero-crossing rate As noted earlier, zero-crossing rate (ZCR) is a simple measure of fricative intensity. It is the number of times points in the discrete-time signal change

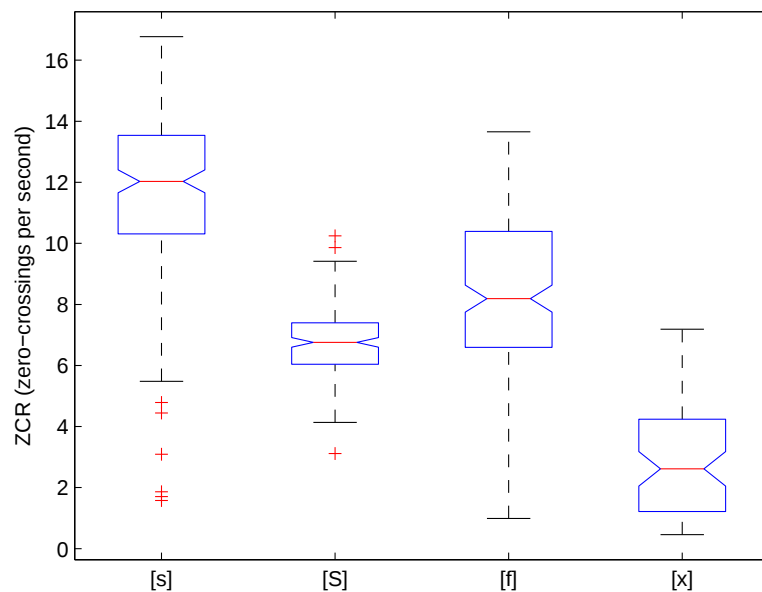


Figure 3.9: Boxplots of Zero-crossing rate (ZCR) by fricative (place of articulation). $F(3, 609) = 297.9$, $p < 0.001$. Tukey’s HSD establishes that all places of articulation are significantly different from each other according to this measure.

algebraic signs in one second (see Section 2.8.3). It appears to be too simple a measure to reliably capture the difference between nasal and oral fricatives. ZCR performed well in differentiating V1 vowel quality ($F(4, 608) = 7.87$, $p < 0.001$), fricative place of articulation ($F(3, 609) = 297.9$, $p < 0.001$), and V2 vowel quality ($F(4, 608) = 6.42$, $p < 0.001$), but it was not useful in distinguishing nasal and oral articulations ($p > 0.05$). The boxplot for fricative place of articulation is presented in Figure 3.9. Tukey’s HSD indicated that all places of articulation are significantly distinct from one another in terms of ZCR.

High frequency spectral energy

HiSlope For an explanation of this measure, see Section 2.8.5. When a single measure of HiSlope is taken across the entire fricative, the resulting measures are unable to distinguish between nasality condition for any speaker. However, the first frame of HiSlope is able to distinguish between V1 produced in a nasal or oral environment ($F(1, 611) = 5.54$,

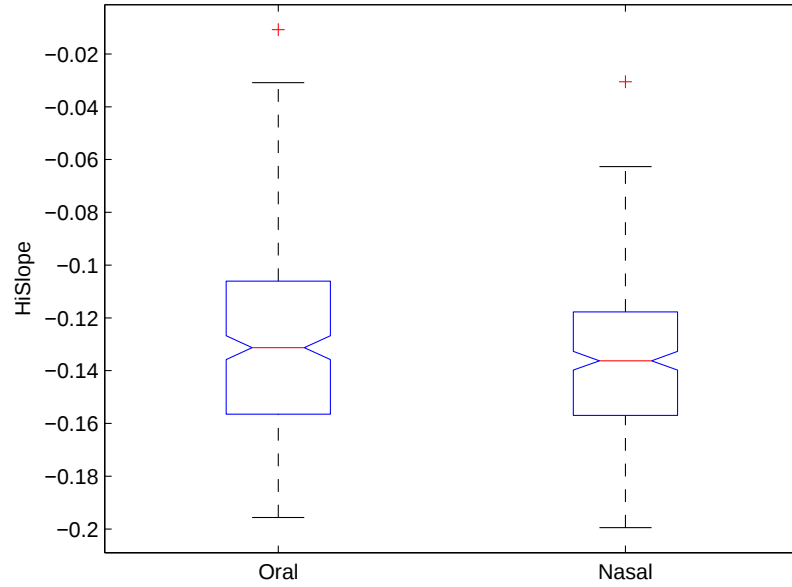


Figure 3.10: Boxplots of HiSlope (dB/kHz) by nasality condition. $F(1, 611) = 5.54$, $p < 0.05$.

$p < 0.05$). Slope under the oral condition is greater (i.e. less negative), suggesting more high frequency energy for the oral fricatives. This effect only obtains in the first few milliseconds after nasalized V1. Thus, here is one indication of the changes introduced by nasalization during fricative production: high frequency spectral energy in some cases declines.

HiBand This is the average spectral energy in a high frequency region of the spectrum, viz. 4–6 kHz (see Section 2.8.5). HiBand measures can be used to successfully distinguish V1 ($F(4, 608) = 21.76$, $p < 0.01$), fricative place of articulation ($F(3, 609) = 404.45$, $p < 0.01$), and V2 ($F(4, 608) = 20.17$, $p < 0.01$) but not (generally speaking) whether the fricative was produced in a nasal or oral context. One exception is for French Speaker 3, where HiBand in the first frame of the fricative is significantly different under these two conditions ($F(1, 142) = 4.36$, $p < 0.05$). The results are presented using boxplots in Figure 3.11.

HiBand in the second ($F(1, 106) = 4.57$, $p < 0.05$), fourth ($F(1, 106) = 7.46$, $p < 0.01$), and fifth ($F(1, 106) = 9.74$, $p < 0.01$) frames for Hindi Speaker 1 distinguish between the nasal and oral conditions of V2.

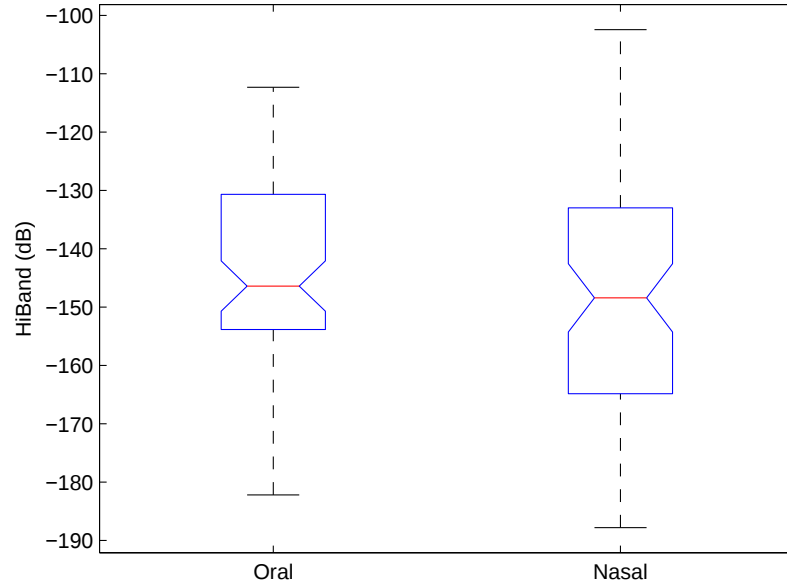


Figure 3.11: Boxplots of HiBand measures (4–6 kHz) in the first fricative frame by nasality condition of V1. $F(1, 142) = 3.36$, $p < 0.05$.

These results more strongly indicate the negative effect of nasality on high frequency energy than the results obtained using HiSlope.

Low frequency spectral energy

LoSlope For an explanation of this measure, see Section 2.8.5. While LoSlope was an effective predictor of fricative place of articulation ($F(3, 609) = 13.91$, $p < 0.001$), it was not effective in discriminating between nasality conditions for any speakers ($p > 0.05$).

Dynamic amplitude For an explanation of the Dynamic Amplitude measure, see Section 2.8.5. For all speakers combined, this variable was a significant predictor of V1, C, and V2 ($p < 0.001$) but not of the nasality of either V1 or V2. This generalization holds true for all of the individual speakers, as well.

Spectral peak bandwidth For an explanation of this measure, see Section 2.8.5. Spectral peak bandwidth, generally speaking, did not perform well in distinguishing fricatives produced under different nasality conditions. One exception is for Hindi Speaker

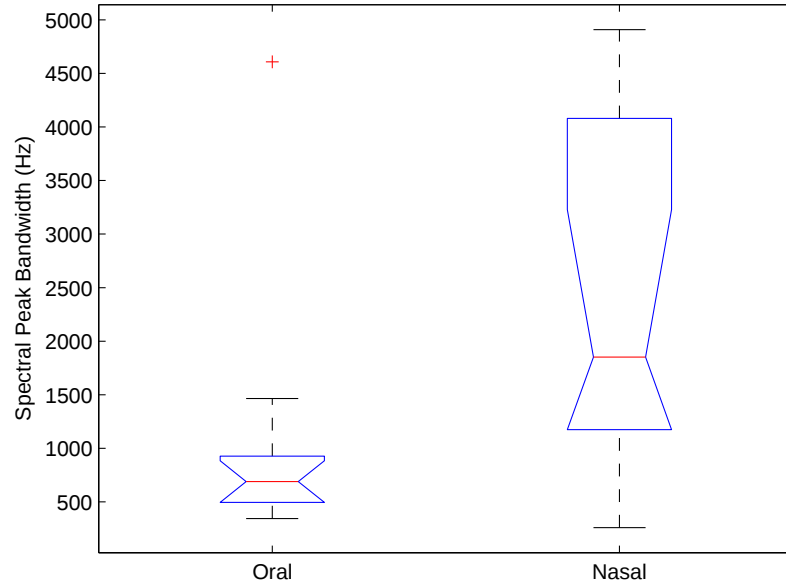


Figure 3.12: Boxplots of spectral peak bandwidth by nasality condition of V1 for speaker A. G. (Hindi). $F(1, 611) = 5.2$, $p < 0.05$.

1 ($F(1, 611) = 5.2$, $p < 0.05$). Results for this speaker are presented in Figure 3.12. They indicate that the nasality of V1 has a considerable effect on spectral peak bandwidth of an adjoining fricative, making the peak a good deal wider (by approximately 1 kHz) than the spectral peak bandwidth of comparable oral fricatives. Furthermore, as can be seen in 3.12, there is a greater degree of variation in bandwidth under the nasal condition.

3.3 Mechanical fricatives

3.3.1 Aerodynamic results

When air was discharged into the fricative model and simultaneously evacuated through pseudo-velopharyngeal ports of increasing size, the pressure in the system dropped. This was the expected result but was still confirmed empirically. The relationship between VPO (in cm^2) and pressure (cm H_2O) for model [s] are reported in Figure 3.13. For smaller VPO, the pressure decrement is of a smaller magnitude than it is for larger VPO. The fact that the pressure for $\text{VPO} = 0 \text{ cm}^2$ is smaller than it is for $\text{VPO} = 0.005$, 0.02 , and 0.045

cm^2 is somewhat puzzling, but the difference is within approximately 1.5 cm^2 . Nonetheless, the correlation coefficient between the two is $r = -0.954, p < 0.001$.

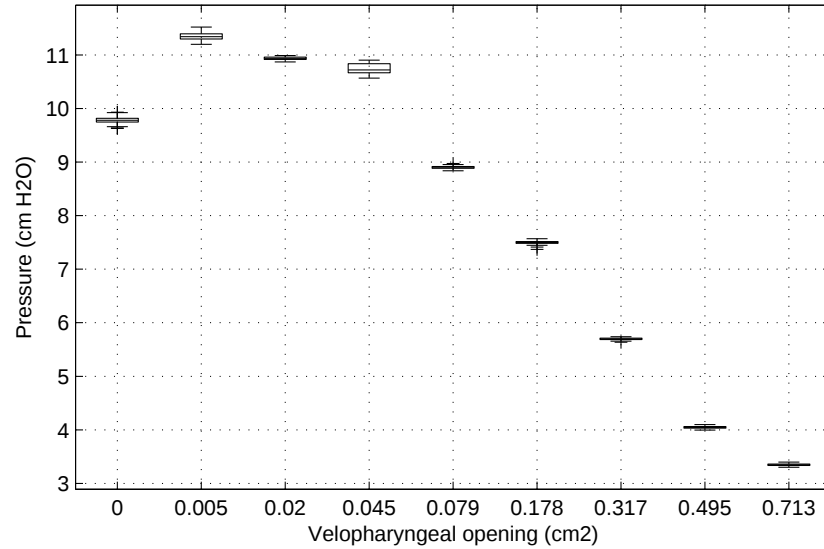


Figure 3.13: The relationship between pressure and pseudo-velopharyngeal aperture during the fricative model [s]. $r = -0.954, p < 0.001$.

To test the significance of differences in pressure between VPO increments, each signal during a given VPO was divided into 100 contiguous samples (5 ms each), and the average pressure value was counted as a trial. Thus, for each VPO, 100 values were used for statistical purposes. The distributions of the samples were not normal according to Lilliefors (see Section 2.11.3), so Kruskal-Wallis (see Section 2.11.4) was used instead of ANOVA. The results showed significant differences in pressure between the various VPO sizes: $\chi^2(8, 891) = 887.77, p < 0.001$. Tukey's honest significant differences were also calculated; the results are reported in Table 3.5.

The results of Tukey's honestly significant differences (along with the correlation coefficient, $r = -0.954$ where $p < 0.001$) indicate that generally speaking an increase in VPO resulted in a decrease in pressure. The relatively minor inconsistencies at the lower end of the VPO range (e.g. $\text{VPO} = 0 \text{ cm}^2$ does not differ significantly from $\text{VPO} = 0.045 \text{ cm}^2$) are unexplained and should be taken into account when comparing the acoustic parameters of fricative noise in this range. In other words, because of the pressure facts, it is safer to draw conclusions based on comparisons of large and small VPO rather than degrees of VPO

Table 3.5: Tukey’s honestly significant differences for pressure by VPO. Groups whose mean is significantly different from a corresponding group ($p < 0.001$) are marked by ‘***’.

	Velopharyngeal opening (VPO) (cm ²)								
	0	0.005	0.020	0.045	0.079	0.178	0.317	0.495	0.713
0	—								
0.005	***	—							
0.020	***		—						
0.045		***		—					
0.079		***	***	***	—				
0.178	***	***	***	***		—			
0.317	***	***	***	***	***		—		
0.495	***	***	***	***	***	***		—	
0.713	***	***	***	***	***	***	***		—

in the small range (e.g. 0–0.045 cm²).

3.3.2 Acoustic results

The mechanical fricatives were recorded with the following nine degrees of velopharyngeal opening (VPO): 0, 0.005, 0.020, 0.045, 0.079, 0.178, 0.317, 0.495, and 0.713 cm² (see Section 2.7.1). Accordingly, nine different values for each acoustic parameter (e.g. zero-crossing rate) are reported.

Each fricative was 1,000 ms long and was analyzed according to the procedures set forth in Jesus and Shadle (2002) and reviewed in Section 2.8.4. Because there should be no ‘coarticulatory’ effects during the production of the mechanical fricatives, ensemble-averaging was not necessary, so the main results are of time-averaged data.³

The time-averaged spectra for the model fricative [s] are presented in Figure 3.14. It should be noted that the high frequency peaks in Figure the bottom panel of Figure 3.14 (the greater VPO condition) are lesser in amplitude than the peaks in the top panel (the 0-VPO condition). Though the 0.713 cm²-VPO peaks seem more prominent, this is relative to the rest of the signal, which on the whole has much less energy than the 0-VPO token. Empirically, there is greater high frequency energy in the token with lesser VPO. This can be seen simply by comparing the data, for example, between 6 and 8 kHz in the two figures.

³However, it seems useful to compare the results of ensemble-averaging on the mechanical and spoken fricatives if only to judge the reliability of the technique where it is more appropriate, i.e. the spoken fricatives. These results are presented after the main results in Section 3.3.2.

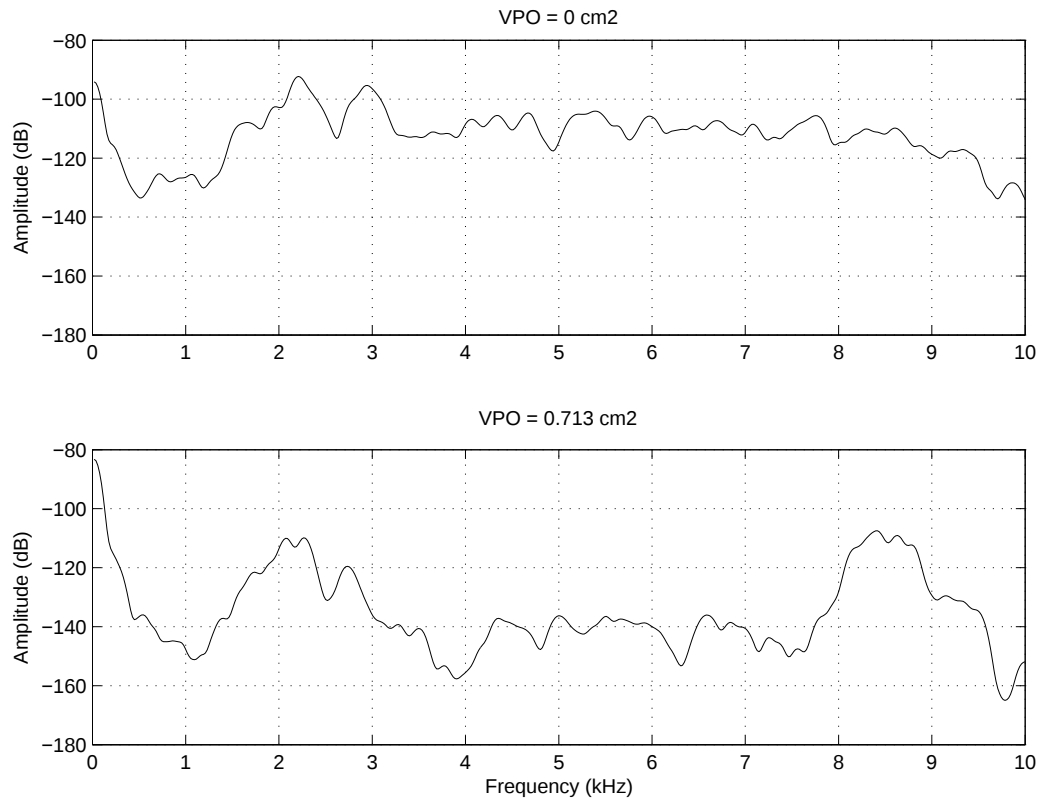


Figure 3.14: The averaged spectra (21 windows, 1024-pt FFT) of mechanical [s] produced with no velopharyngeal opening ($VPO = 0 \text{ cm}^2$) (top panel) and with $VPO = 0.713 \text{ cm}^2$ (bottom panel).

$\bar{F}$ by velopharyngeal opening Results for $\bar{F}$ of the mechanical fricatives by velopharyngeal opening (in kHz) are shown in Table 3.6. The correlation coefficient, 0.344 fails to achieve significance even at $p < 0.05$. No clear pattern emerges. The results are presented graphically in Figure 3.15.

Zero-crossing rate Results for zero-crossing rate, defined in Section 2.8.3, are presented in Table 3.7. According to Rabiner and Schafer (1978: 127), “[A] zero-crossing is said to occur if successive samples [in a discrete-time signal] have different algebraic signs.”

High frequency spectral energy

Table 3.6: $\bar{F}$ of mechanical fricatives (in kHz) at differing velopharyngeal openings

FRIC	Velopharyngeal opening (VPO) (cm ²)									r
	0	0.005	0.020	0.045	0.079	0.178	0.317	0.495	0.713	
s	3.11	2.27	2.94	2.52	2.94	4.62	2.52	2.36	4.29	0.344

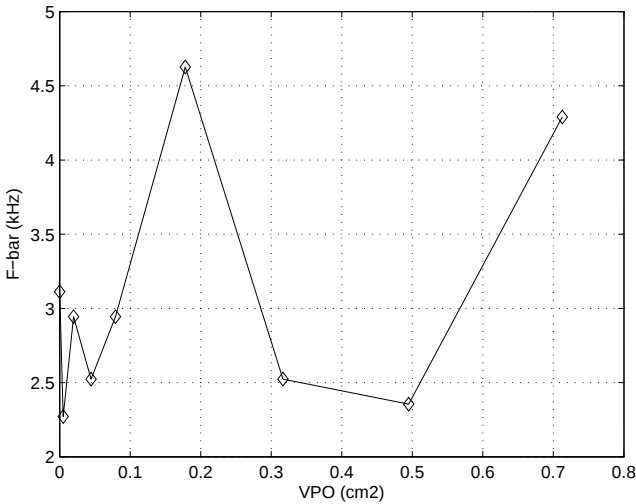


Figure 3.15: $\bar{F}$ (kHz) measurements for mechanical [s] produced at a range of velopharyngeal openings.

HiSlope Results for HiSlope, defined in Section 2.8.5, are presented in Table 3.8. The results are presented graphically in Figure 3.16.

HiBand Results for HiBand, defined in Section 2.8.5, are presented in Table 3.9. The results are presented graphically in Figure 3.17. HiBand measures for the mechanical friative [s] suggest a strong, inverse correlation ($r = -0.989$, $p < 0.001$) between this variable and velo-pharyngeal opening, i.e. as velopharyngeal opening increases, high frequency energy decreases. A similar effect was observed in some of the recorded data,

Table 3.7: Zero-crossing rate of mechanical fricatives at differing velopharyngeal openings

FRIC	Velopharyngeal opening (VPO) (cm ²)									r
	0	0.005	0.020	0.045	0.079	0.178	0.317	0.495	0.713	
s	8.7	9.325	6.7437	6.1147	5.4085	5.2438	5.6398	5.7353	7.5101	-0.210

Table 3.8: HiSlope (dB/kHz) of mechanical fricatives at differing velopharyngeal openings

FRIC	Velopharyngeal opening (VPO) (cm ²)									<i>r</i>
	0	0.005	0.020	0.045	0.079	0.178	0.317	0.495	0.713	
s	-0.015	-0.003	-0.018	-0.020	-0.019	-0.029	-0.010	0.001	0.031	0.800*

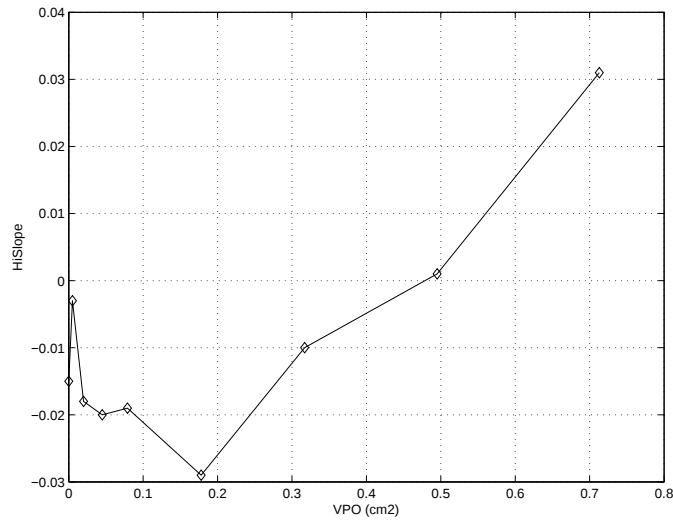


Figure 3.16: HiSlope (dB/kHz) measurements for mechanical [s] produced at a range of velopharyngeal openings.

e.g. for Speaker 3 (see Figure 3.11, though the effect was not robust across speakers as the model data suggest).

Table 3.9: HiBand of mechanical fricatives at differing velopharyngeal openings, $r = -0.989$, ($p < 0.001$).

FRIC	Velopharyngeal opening (VPO) (cm ²)								
	0	0.005	0.020	0.045	0.079	0.178	0.317	0.495	0.713
s	-108.5	-102.1	-100.9	-105.8	-112.7	-117.7	-128.9	-141.0	-156.4

Low frequency spectral energy

LoSlope Results for LoSlope, defined in Section 2.8.5, are presented in Table 3.10. The results are presented graphically in Figure 3.18. LoSlope was not an illuminating variable in the discrimination of nasal and oral fricatives for the spoken data. LoSlope is

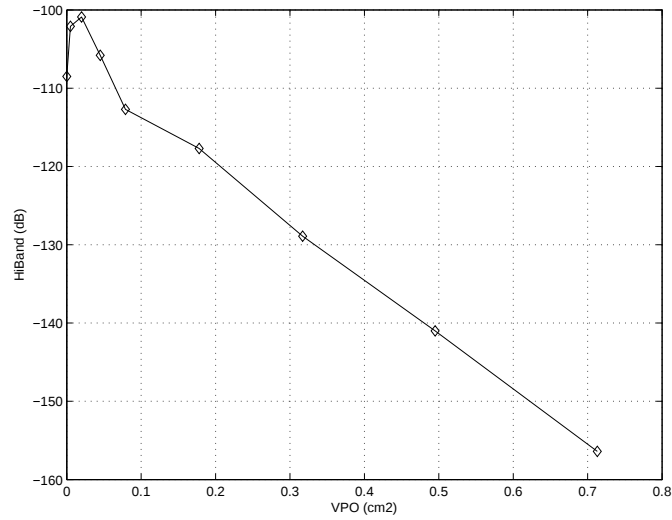


Figure 3.17: HiBand measurements for mechanical [s] produced at a range of velopharyngeal openings.

only marginally correlated with VPO for mechanical [s] ($r = -0.546$, but the correlation is significant $p < 0.05$).

Table 3.10: LoSlope (dB/kHz) of mechanical fricatives at differing velopharyngeal openings

FRIC	Velopharyngeal opening (VPO) (cm ²)									<i>r</i>
	0	0.005	0.020	0.045	0.079	0.178	0.317	0.495	0.713	
s	0.880	0.107	0.390	0.733	0.0360	-0.014	0.188	0.197	-0.546	-0.699*

DynAmp Results for Dynamic Amplitude, defined in Section 2.8.5 are give in Table 3.11. The results are presented graphically in Figure 3.19. This measure is strongly and postively correlated with VPO. These results are strikingly incongruous with those from natural speech, where neither the nasality of V1 or V2 could be reliably predicted based on DynAmp.

Table 3.11: DynAmp (dB) of mechanical fricatives at differing velopharyngeal openings, $r = 0.961$, $p < 0.001$.

FRIC	Velopharyngeal opening (VPO) (cm ²)								
	0	0.005	0.020	0.045	0.079	0.178	0.317	0.495	0.713
s	5.134	2.686	4.421	4.330	8.686	9.554	17.643	19.783	22.4337

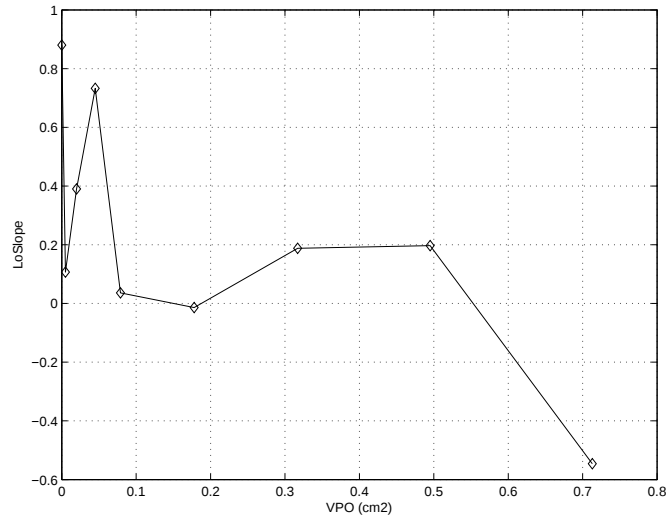


Figure 3.18: LoSlope (dB/kHz) measurements for mechanical [s] produced at a range of velopharyngeal openings.

Spectral peak bandwidth Results for Spectral peak bandwidth, defined in Section 2.8.5, are presented in Table 3.12. The results are presented graphically in Figure 3.20. For the spoken data spectral peak bandwidth was a strong predictor of nasality for only one speaker (see Section 3.12).

Table 3.12: Spectral peak bandwidth of mechanical fricatives at differing velopharyngeal openings (in Hz)

	Velopharyngeal opening (VPO) (cm ²)									
FRIC	0	0.005	0.020	0.045	0.079	0.178	0.317	0.495	0.713	<i>r</i>
s	78	58	59	117	236	253	253	410	391	0.9065***

Ensemble-averaged data for mechanical fricatives

Ensemble-averaged data consists of acoustic measures taken from individual windows in a fricative then averaged together to represent the time-varying aspects of the noise in a frame-by-frame analysis. Because ‘coarticulatory’ variation was assumed to be minimal during the mechanical fricatives, it is not necessary to conduct a rigorous analysis using ensemble-averaged data. Nevertheless, ensemble-averaged data from the mechanical fricatives may still be put to good use, e.g. as a point of comparison with the spoken fricatives which were indeed coarticulated. One assumption of the methodology in Jesus and Shadle

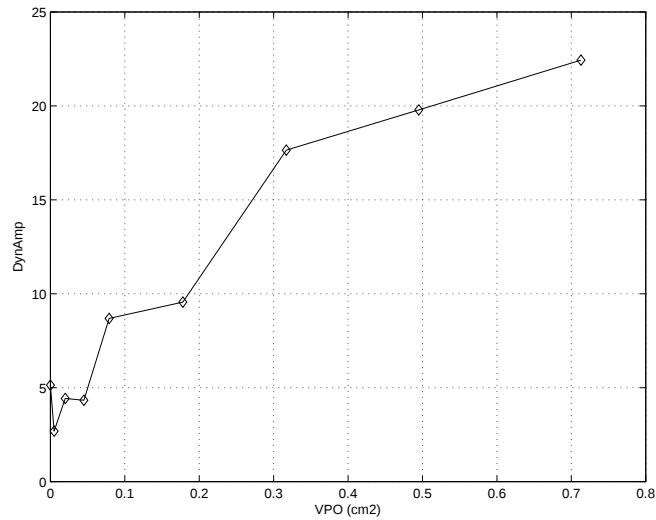


Figure 3.19: DynAmp (dB) measurements for mechanical [s] produced at a range of velopharyngeal openings.

(2002) is that there is significant variation between individual portions of a fricative due to coarticulation. Without the coarticulation, there should be no significant frame-by-frame differences. The mechanical fricatives provide a test case. Inter-frame variation for the mechanical fricatives can be compared with inter-frame variation for the spoken fricatives. Thus, the standard deviations of the various acoustic measures are reported in Table 3.13 below for the mechanical fricative [s].

Table 3.13: Frame-by-frame variation in mechanical fricatives ([s] at 9 degrees of velopharyngeal aperture) for four different acoustic measurements. For example, the standard deviation in HiBand for 21 frames of a mechanical fricative produced at 0 cm² VPO is 3.50 dB.

	Velopharyngeal opening (VPO) (cm ²)								
	0	0.005	0.020	0.045	0.079	0.178	0.317	0.495	0.713
HiSlope	0.01	0.07	0.01	0.01	0.03	0.01	0.02	0.10	0.02
LoSlope	1.28	2.33	1.43	1.65	2.42	2.50	2.74	3.41	3.72
DynAmp	15.98	15.07	13.11	10.13	14.47	11.82	12.69	13.88	11.23
HiBand	3.50	2.71	3.92	3.91	2.58	3.24	2.73	3.03	3.08

There is no discernible pattern in Table 3.13. It may be sufficient to observe that frame-by-frame variation for this mechanical fricative is not very different at differing VPO values. This stands out in contrast to the results presented in Section 3.2.2, where an effect

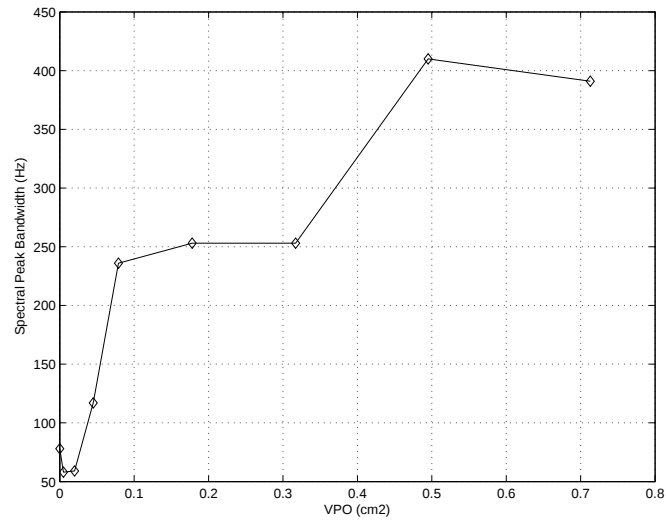


Figure 3.20: Spectral peak bandwidth measurements (Hz) for mechanical [s] produced at a range of velopharyngeal openings.

sometimes achieved significance only for a certain frame (e.g. for the first frame of HiSlope in predicting the nasality of V1, see Figure 3.11).

Chapter 4

Discussion and Conclusions

4.1 Summary of the results

The results of this study highlight an important finding in the ongoing nasalized fricative controversy: fricatives can be nasalized, which leads to the modification of certain spectral properties. The acoustic effects of nasalization on spoken voiceless fricatives have been carefully examined in the present study, but they do not lead to firm conclusions about the acoustic debilitation of fricatives in nasalized contexts.

A significant difficulty was mediated—though not entirely overcome—in this thesis. The simultaneous recording of aerodynamic and high-quality acoustic signals (using the conventional mask methodology) is highly problematic.¹ As demonstrated in Figures 2.2 and 2.3, the utilization of aerodynamic and acoustic methods in tandem sometimes has unforeseen repercussions on the data. It is therefore not impossible that in the present study, some acoustic tokens have been counted as ‘nasalized’ when in fact there was not a significant degree of nasal airflow at the time of their utterance. While the author recategorized tokens that sounded less nasal than the stimuli presented to the speaker, e.g. when the speaker mistook Brazilian Portuguese *arrã* for [axa], this cannot be considered a fullproof method. Moreover, even assuming no errors in the pronunciation of vowels as ‘nasal’ or ‘oral’ among the tokens, nothing can be said of the relative degree of their nasalization.

Here, data from the mechanical fricatives at least partially filled the lacuna. Be-

¹Hot wire anemometry or pneumotachography seem like suitable supplements, if not replacements, to the mask methodology. However, they could not be attempted in the present time frame (Cotes et al. 2006: 62–63).

cause the pseudo-velopharyngeal aperture of the model fricative could be adjusted to mimic varying degrees of aperture in an actual vocal tract, the problem of gradient nasalization could be dealt with, though only indirectly.

Despite this versatility, however, the model data can only approximate what is occurring in an actual vocal tract. The differences between the model data and the spoken data seem great enough to engender skepticism as to whether or not one is really a reflection of the other. For example, the effect of spectral peak bandwidth seems extremely relevant in the model data (see Figure 3.20), with greater velopharyngeal aperture increasing the measure significantly. However, among the spoken data, the same effect was found for only one speaker.

There are several possible reasons for the discrepancy. Perhaps the effects of coarticulatory nasalization on fricatives are so small that many more subjects are needed to bring them into sharper focus. The time-consuming nature of performing aerodynamic recordings and the physical awkwardness (if not discomfort) of the procedure placed severe limits on the number of subjects that could be included in the present study. Only future studies can contemplate a larger speaker base. In any case, until a strategy can be developed to capture a high-quality acoustic signal (one amenable to the detailed acoustic parameterization presented by Jesus and Shadle (2002)) and an accurate aerodynamic signal, the present conclusions are only tentative ones.

Another, perhaps more interesting, possibility is that speakers may be able to compensate for the deleterious effects of nasalization by increasing airflow. With the velum lowered, it is possible that speakers routinely make adjustments in transglottal flow just enough to overcome the velopharyngeal escape and maintain the acoustics of the fricative. Indeed, one may presume that speakers with relatively minor velopharyngeal dysfunction do this as a matter of course. For speakers with major velopharyngeal dysfunction, it has been shown that the acoustics of fricatives are drastically altered (Weinberg and Horii 1975). If this view is taken, then the model data are extremely relevant, in that they present us with a picture of a system that lacks a compensatory feedback loop.

How one might demonstrate the existence of compensatory transglottal flow during a nasalized fricative is not altogether clear, since any compensation made ‘upstream’ of the velopharyngeal opening would be depleted (through the nose) before the flow reached any external recording device. Furthermore, depending on the degree of velopharyngeal opening, the difference may be quite small. Plethysmographic evidence might help settle

the question, as the activity of the lungs during nasalized and oral fricatives could be shown to differ significantly under the two conditions.

In sum, the present study contemplates the acoustic features of fricatives that may be modified by the presence of an open velopharyngeal port (i.e. nasalization), thus inhibiting the phonologization of nasalized fricatives. The high frequency energy of fricatives and their narrow spectral peak bandwidth are likely to fall victim to nasality. The importance of high frequency energy in the production and perception of fricatives is well known (Johnson 1997, Stevens 1998, Jesus and Shadle 2002). Narrow spectral peak bandwidth, on the other hand, is not discussed as widely in the fricative literature. It has been dealt with, so far as I am aware, in only one study, and that is of the relatively uncommon ‘whistled fricatives’ [ɕ ʑ] of Shona² (Bladon et al. 1987). Whether spectral peak bandwidth is a measure useful in perceptually differentiating [s] from [f], for example, remains to be seen.³ If these variables are indeed essential to fricative perception, then the alteration of their values under the effects of nasalization may be considered disruptive to an otherwise orderly phonemic inventory.

Based on the present results, I conclude that voiceless nasalized fricatives like [ɕ̃ ʑ̃] may occur epiphenomenally in the languages of the world, but not without significant changes to their spectral characteristics. The prediction of spectral change is based on a constant regime of airflow rather than one in which transglottal airflow subtly increases during the fricative. In cases of compensatory airflow, the increase could make up for any nasal escape, especially at low levels of VPO, resulting in a relatively unaltered fricative spectrum.

I further conclude that it is not unreasonable to posit nasal harmony systems that allow for the lowered velum during the production of fricative sounds (such as Coatzospan Mixtec), with the following caveat: The language in question should not allow nasalization to occur through ‘peaked’ fricatives like [s ʃ] if the language already has flat-spectrum fricatives like [f x θ]. As evidenced by the model data, nasalization of [s] could widen its spectral peak bandwidth and reduce its high frequency energy, causing it to become more like a flat-spectrum fricative. If such fricatives already exist in the language, it would be difficult

²Whistled fricatives in Tshwa (Tshwa-Ronga, Mozambique) have been discussed by Shosted (2006b).

³The notion of an acoustic-perceptual space for fricatives has traditionally received less attention than the notion of vowel space. The reason for this is straightforward: The parameterization of vowels using F1, F2, and F3 makes study of the vowel space possible, while parameterization of fricative space has not, so far, been successful. Nevertheless, some type of ‘fricative space’ must exist in languages with multiple fricative phonemes.

to distinguish, e.g. [š] from [f]. This predicts nasal harmony systems *unlike* Applecross Scots Gaelic, where there are numerous flat-spectrum fricatives and peaked fricatives, all of which may undergo nasalization. By the same reasoning, flat-spectrum fricatives are unlikely to undergo phonemic nasalization regardless of the number of fricatives, since it seems unlikely that nasalization would significantly alter their acoustic signatures.

Model data in the present study clearly demonstrate that the degree of velopharyngeal opening plays an important role, as spectral characteristics such as high frequency energy and spectral peak bandwidth are significantly altered only as the velopharyngeal port opens more widely. Thus, nasalization during fricatives must be seen as a gradient phenomenon. While it may occur at relatively low levels with no severe acoustic cost, the same cannot be expected as VPO increases.

These findings have implications for a wide variety of geographically and typologically diverse languages said to have voiceless nasalized fricatives (see Section 1.7). It suggests that the perceptual salience of voiceless nasalized fricatives is weakened and that they are more likely to be confused with fricatives at other places of articulation. For example, [š] may be confused with [x] because both have relatively low-amplitude energy in the high frequencies and broad peak bandwidths. On the other hand, a fricative like [χ̃] may not be adversely affected by nasalization. Thus, fricatives with relatively flat spectra (e.g. [f x θ]) are more likely to be epiphenomenally nasalized than fricatives with large spectral prominences (e.g. [ʃ s]). In a language without oral flat-spectrum fricatives, [š̃ ʃ] could reasonably stand in phonemic opposition to [s ʃ].

While such phonological patterns may be posited based on present experimental data, they do not happen to appear in the languages of the world in which nasalized fricatives are claimed to exist. Moreover, they do not appear influential in nasal harmony systems in which nasality is allowed to ‘spread’ through fricatives (see Section 1.7). If [š] is just as common as [f̃], for example, the compensatory transglottal flow hypothesis might be invoked. To wit, we can assume from the spectral characteristics of [s] and the findings of the present study that the acoustics of [s] are more likely to be impaired by nasalization than the acoustics of [f]. If, however, transglottal flow is increased, just for the articulation of [š̃], then there is no reason to believe that it cannot occur as often as [f̃], which, unimpaired by the open velopharyngeal port, requires no compensatory flow. As can be seen, much rests on the further elaboration and testing of the compensatory flow hypothesis in order to straighten out these claims.

No language of the world has a voiceless, buccal, nasalized fricative that occurs phonemically. The findings of the present study do not, however, rule this out as a possibility.

4.2 Nasal harmony

In her thesis on nasal harmony, Walker addresses the issue of consonants that either ‘block’ or allow nasalization to ‘spread’ throughout a prosodic constituent (2000) (see Section 1.7.11). From the Ohalian point of view (at least the strong hypothesis—see Section 1.6), fricatives pose an obstacle to a coarticulatory account of nasal harmony. Imagine a language in which the segment [n] triggers rightward-spreading nasalization throughout the entire word. In a form like /nəsi/ the expected outcome would be [nẽsi]. What occurs during the [s]? According to the strong version of the Ohalian hypothesis, the fricative may not be nasalized, so the erstwhile lowered velum has raised to allow the full production of the alveolar fricative. Afterwards, it lowers again during the production of [i]. Coarticulation (or at least phonetic ‘coproduction’, i.e. gestural overlap in the sense of Browman and Goldstein (1986)) cannot account for the nasalization of the last vowel, since the velum is lowered on two separate occasions. Whatever motivates the nasal harmony, one cannot argue that it is coarticulation. Unless, of course, the /s/ is realized as [š], countering the strong version of the hypothesis. According to the weaker version of the Ohalian hypothesis (see Section 1.6), [š] may occur phonetically but it cannot achieve the status of a phoneme. It would seem that coarticulation can explain nasal harmony that acts through fricatives as long as /s/ is not contrastive with /š/. However, the results of the present study suggest that [š] and [ʃ̃] may be acoustically more similar to fricatives like [f] and [x], complicating the matter of nasalization ‘spreading’ equally through all fricatives.

Walker mentions 28 languages in which all segments (including fricatives) allow nasalization to ‘spread’ (2000). These are listed, along with the complete fricative inventories of 24 of the languages, in Table 1.9. Of these 24 languages, the average number of voiceless fricatives per language is approximately 2.5. Half of the languages have an opposition between a flat-spectrum fricative like [f] or [x] and a sibilant phoneme like [s] or [ʃ].⁴ This typological evidence is not exactly what we would expect based on the confusability of,

⁴Because there is no aerodynamic reason to believe [h] cannot be nasalized (as mentioned in Section 1.1), the glottal fricative is not counted as one of these ‘flat-spectrum’ fricatives.

e.g. [š] and [x], suggested by the acoustic experiments conducted here.⁵ Accordingly, some of the most enlightening phonetic information regarding phonetically nasalized fricatives might come from Northern and Southern Cabécar (Chibchan, Costa Rica), Epera (Choco, Panama), Gbeya (Niger-Congo, Central African Republic), Gokana (Niger-Congo, Nigeria), and Guaraní (Tupí, Paraguay), languages with ‘peaked’ and flat-spectrum fricatives and nasal harmony that acts through both.

Of these languages, Guaraní undoubtedly has the largest speaker population and should perhaps be the first to undergo a serious investigation. The aeroacoustics of Guaraní [s ʃ x] in nasal and oral domains would test the results of the present study. How dissimilar are Guaraní [š] and [s]? How similar are [š] and [x]?

Walker’s typological data suggest that languages are not always constrained according to my predictions. In other words, nasalization of sibilant fricatives may occur in languages that have flat-spectrum fricatives (if we assume that segments that allow the ‘spread’ of nasality are nasalized in the process). Nonetheless, based on the present results, it seems more plausible that a language like Tucano, with only the fricatives [s h], should allow these to nasalize phonetically because [š] and [h] are acoustically dissimilar. On the other hand, a language like Applecross Scots Gaelic (see Section 1.7.1), with a total of 6 voiceless fricatives—all of which may be phonetically nasalized—stretches the imagination. How could [š ɬ ʃ ɸ ɰ ɸ̥] possibly be distinguished from one another (and their oral counterparts) if their spectral properties are altered as the present study suggests?

In sum, this dissertation elaborates and makes predictions about the role of aeroacoustics in nasal harmony, predicting that sibilant fricatives are most likely to block nasalization because they have the most to suffer acoustically.

4.3 Velopharyngeal dysfunction

It has been shown that in the speech of individuals with velopharyngeal dysfunction (i.e. cleft palate), [š] is spectrally similar to a velar or pharyngeal fricative (Weinberg and Horii 1975). This observation is also supported by the present study, insofar as the decrease of high frequency spectral energy and spectral peak bandwidths may be said to figure prominently in the production of velar and pharyngeal fricatives (Jesus and Shadle 2002).

⁵As stated before, this confusability is at present suppositional, pending further perceptual work on the acoustic variables in question.

Weinberg and Horii found that a consistent feature of cleft palate speakers' /s/ was the presence of multiple spectral maxima. Furthermore, they concluded that low frequency excitation of F2 in the /s/ of cleft palate speakers was generally comparable to low frequency excitation in Arabic /ħ/. As Weinberg and Horii note, cleft palate speakers often make articulatory adjustments in the production of fricatives, moving the place of greatest constriction upstream of the velopharyngeal port. The speakers' adaptation overcomes the aerodynamic problem of the 'leaking valve' by removing the place of articulation to a point upstream of the leak.

However, it is also possible that no articulatory adjustment was made and that Weinberg and Horii's data are records of [ɕ]⁶ rather than [ħ]. The authors do not address the controversy of nasalized fricatives. Their interest was primarily in the acoustics of the sound produced, not the physiological adaptation that may or may not have been effected by their subjects. Whether a pharyngeal constriction was or was not made can only be surmised. The disruption of the spectrum, however, may be easily attributed to the presence of nasalization. Further research in this area, with appropriate controls for actual place of articulation, are warranted.

4.4 Voiceless nasals

At this point, it may be advantageous to distinguish the relatively well-studied class of sounds known as 'voiceless nasals' from nasalized fricatives. In various languages of Southeast Asia, including Burmese, Hmong (Hmong-Mien, Thailand), and Iaa (Austronesian, New Caledonia), a set of voiceless nasals like [m̥ n̥ ŋ̥] stand in contrast to modal-voiced nasals like [m n ŋ]. With a wide-open glottis and air rushing through the nostrils, the closed oral 'sidebranch' of the system contributes relatively little to the acoustic output of a voiceless nasal, thus making the fricative portion of the various voiceless nasals relatively difficult to differentiate from one another. Based on this reasoning, as well as recordings of the sounds (Ladefoged and Maddieson 1996), it is generally agreed, as observed by Ladefoged (1971) and Ohala (1975), that voicing at the offset of the consonant is helpful in distinguishing place of articulation among voiceless nasals. Thus /m/ is routinely realized [m̥m], etc., and the cues for place of articulation are to be found in the acoustic material during the voiced portion of the sound.

⁶Speech pathologists might prefer the transcriptions [ɕ̥] 'nasal escape' or [ɕ̥̥] 'velopharyngeal friction'.

Nasalized fricatives are, physiologically speaking, a much different subject. Traditional usage of the term ‘nasalized’ rather than ‘nasal’ implies a secondary articulation. Thus, we take it for granted that the primary articulation of [ʃ̃], for example, occurs at the alveolar ridge, not at the nostrils, as is the case for [ɲ̃]. The nomenclature implies that for nasalized fricatives, the dominant airflow is oral, with some nasal airflow, whereas during a voiceless nasal, the dominant airflow is nasal. As with Gerfen (1999, 2001), the present study suggests that nasal airflow may occur at least at the edges of a fricative in a nasal environment (see Section 3.2.1), substantiating the phonetic existence of segments like [ʃ̃] and casting doubt on the strong version of the Ohalian hypothesis (see Section 1.6).

Despite their differences, one wonders if there is not some relation between voiceless nasals, whose existence is undisputed, and voiceless nasalized fricatives, which are more controversial. Is it reasonable to propose that in the diachronic development of voiceless nasals they passed through a stage as nasalized fricatives? For example, in Burmese, the historical form /sn/ is realized in the modern language as /ɲ̃/, e.g. Written Tibetan *sna* → Burmese /ɲ̃a/ ‘nose’ (Greenlee and Ohala 1980).⁷ It seems there are two possible reasons for this change.

The first explanation, perhaps the more obvious, is that of perseverative assimilation: the vocal folds, spread wide during the articulation of voiceless /s/ do not achieve modal vibration (i.e. voicing) until later in the nasal segment, giving rise to a partially devoiced cluster like [sɲ̃n]. Over time, the entire cluster is reinterpreted based not on the voicing of [n] but on the voiceless frication of [sɲ̃]. However, since the voiceless nasal is the quieter sound, the final development to a unitary element [ɲ̃] seems to suggest that the *less* salient (and quieter) voiceless nasal is favored by listeners at the cost of the relatively more salient (and louder) /s/.

A regressive explanation for this diachronic change avoids this problem of salience. If the velum is lowered in anticipation of /n/ during the production of /s/, a voiceless nasalized fricative [ʃ̃] will result. The reduction in fricative intensity caused by nasalization makes the acoustic output similar to, and thus reinterpretable as, the characteristically flat spectra of voiceless nasal consonants. One possible explanation for this diachronic change is that the prominent spectral characteristics of the alveolar fricative were ‘flattened’ by the

⁷Similarly, Sturtevant (1940) and Thurneysen (1946) claim that /s/ + resonant clusters became voiceless nasals in Primitive Greek and Old Irish, respectively. According to Saksena (1971: 45), some breathy voiced nasals [ɲ̃fi] in Awadhi derive from old Indo-Aryan /sn/ clusters, as well.

presence of nasalization in the subsequent consonant. Thus, a relatively salient [s] and a relatively quiet [n̥] do not abut one another at the medial stage of the development. The cluster would instead look like [s̥nn] and the relatively more salient [n̥] would dominate in the ear of the listener. Progressive assimilation need not be invoked to device the /n/ entirely, since only the initial portion (adjacent to the /s/) is voiceless. This explanation for the development of voiceless nasals is hypothetical only, and deserves further attention, as do the acoustic properties of voiceless nasals in a variety of languages.⁸

4.5 Sibilants and non-sibilants

The present analysis deals with two classes of fricatives and their different reactions to nasalization. The claim has been made that, while both classes of fricatives will experience the same effects of nasalization, one is less likely to result in a different percept on the part of the listener. Fricatives with a peaked spectrum like [s ʃ] will experience a lowering of high-frequency energy and a widening of the spectral peak bandwidth. Flat-spectrum fricatives such as [f θ x] will experience the same changes, but since these fricatives already have relatively flat spectra and wide bandwidth peaks, it is assumed that [f̃ θ̃ x̃] will not sound much different from their non-nasalized counterparts. Conversely, [s̃ ʃ̃] will bear less resemblance to [s ʃ] precisely because the acoustic alterations involve spectral characteristics unique to sibilant fricatives. This hypothesis is of course informed by the traditional classification of ‘sibilant’ and ‘non-sibilant’ fricatives, which will now receive some attention.

In the Jakobsonian system of phonological features, the label STRIDENT served primarily to differentiate the labiodental fricatives [f v] from the bilabial fricatives [ɸ β] (Ladefoged 2006). The result was a rather unnatural class of fricatives: [f, v, s, z, ʃ, ʒ]. Noting this irregularity, Chomsky and Halle (1968) regarded [f v] as non-strident fricatives, patterning against the others. The label STRIDENT is supplanted by SIBILANT in Ladefoged (2006). He observes that the term ‘sibilant’ was used as early as the 17th century by the phonetician Holder (1669) to identify [s z ʃ ʒ],⁹ as a natural class.

Is there an articulatory definition that distinguishes the sibilant from non-sibilant fricatives? Ladefoged (2006) observes that sibilant sounds are produced with a raised jaw, such that there is a narrow gap between the upper and lower front teeth. He notes that the

⁸Maddieson (1983) suggests differences in the spectra of the ‘fricative’ portions of different voiceless nasals but understandably does not compare the spectra to fricatives made at similar (oral) points of articulation.

⁹Holder (1669) did not actually recognize [ʒ], a more modern development, as a sound of English.

high frequency aperiodic acoustic energy typical of such sounds arises when the jet of air strikes this narrow gap (see Catford (1977), Shadle (1985), and Section 1.2).

Ladefoged (2006) raises two objections to the jaw-raising hypothesis. First, other sounds not typically understood as sibilants are accompanied by considerable jaw raising, e.g. the high front vowel [i]. Second, he observes that there is “no evidence showing that jaw position is a salient characteristic of sounds causing them to be grouped together.” An acoustic-perceptual account of sibilant fricative relatedness is given in the data of Miller and Nicely (1955) and Shepard (1972). Ladefoged (2006) concludes that “the well attested salient auditory characteristics [shared by sibilants] are clearly the basis for the natural class.”

Based on the acoustic-perceptual definition of a sibilant as being characterized by high-frequency aperiodic energy and a narrow peak bandwidth, the present study argues that sibilants have ‘more to lose’ acoustically and perceptually from velopharyngeal venting. While the acoustic changes are the same for non-sibilants and sibilants, due to fundamental differences in their acoustic structure, nasalization would rob sibilants of perceptually unique and unifying characteristics, while simply increasing by some degree the non-sibilant characteristics of the non-sibilants. In other words, the results of the present study lead us to characterize the nasalization of fricatives as a de-sibilantizing process.

4.6 Universals, rarities, and the expanding IPA

Over a decade ago, Ladefoged wrote (somewhat pessimistically, I think) that “[i]t is becoming harder and harder to mine the phonetic dross and come up with something new” (Ladefoged 1990: 70). Despite at least fifty years of phonetics research informed by advanced methods of digital signal processing, there are still many fundamental questions that keep experimental phoneticians and laboratory phonologists engaged in exploring the physical world of human vocal production. While it may indeed be harder to find a speech sound previously undescribed, there are still many questions worth exploring. For example, the case of ‘nasalized fricatives’ brings into focus a number of issues at the core of phonetics and laboratory phonology, among them:

1. What is phonetically impossible and phonetically implausible?
2. What universal characteristics of the anatomical vocal tract help shape phonological

and typological patterns?

3. How good are physical principles (e.g. physiology, aerodynamics, and perception) at constraining the content of sound systems?

In this concluding section of my dissertation, I will address a few ways in which nasalized fricatives fit into the ‘bigger picture’ of phonetics and even formal phonology.

To summarize, Ohala (1975), Ohala and Ohala (1993), Solé (1999), and others reasonably claim that nasalized fricatives cannot exist. Schadeberg (1982), Gerfen (2001, 1999), Lastra (1984), Stringer and Hotz (1973), and Ternes (1989) claim that they do. The present study weighs in somewhere in the middle, assessing the acoustic potential for phonologization among sibilant and non-sibilant nasalized fricatives.

Upon reflection, the problem of nasalized fricatives highlights the following with regard to current thinking in phonetics and phonology:

1. Phonetic universals are best posited upon consideration of physical mechanisms and perceptual outcomes (i.e. “speech perception is hearing sounds, not tongues” (Ohala 1996));
2. The IPA is indeed expanding in fairly unpredictable ways as phoneticians collect more information about a larger number of diverse languages;
3. Our current understanding of the phonetic characteristics that lead to phonemic outcomes is still lacking.
4. We cannot presently conclude that sound systems consist of a discrete formal system with a limited number of phonological “atoms” (elemental features like $[\pm \text{VOICE}]$ or graphic symbols like [h]) (Port and Leary 2005).

4.6.1 An infinite phonetic alphabet?

Conceived over a century ago, the International Phonetic Alphabet (IPA) aims to provide a symbol for every contrastive element in any given human language¹⁰ (MacMahon

¹⁰The IPA in fact falls short of this goal in several significant respects, e.g. dental vs. alveolar and laminal vs. apico-alveolar consonants, as well as long vs. short vowels.

1996). Diacritic marks are used to indicate subphonemic variations. Since nasalized fricatives are nowhere claimed to be phonemic, it is appropriate that they should be symbolized as an oral fricative with a diacritic tilde, e.g. [x̃].

As our understanding of subphonemic variation increases, i.e. as we collect more data about how seemingly similar phonemes are actually articulated in different ways across languages and speakers, we are confronted with an infinitely expandable IPA. To pose an extreme hypothetical: Should there exist a unique diacritic or scalar value in association with every vowel quality produced by every speaker of every known language? Should these values and/or symbols be encoded in transcription? What does the ideal IPA transcription look like? Would an ideal IPA transcription provide enough information for someone to reproduce an utterance exactly as it was first spoken? Surely, the information load would be great, and the law of diminishing returns would set in fairly quickly, as speech recognition engineers understand.¹¹

Thus, there is a fundamental tension in phonetics and phonology between the search for language universals—those components of sound systems that are relatively invariant across languages—and a universal sound system that can be elaborated virtually *ad infinitum*. Indeed, one may wonder at the universality and systematicity of the result. Port and Leary ask and answer their own question: “Do phoneticians generally agree with phonologists that we will eventually arrive at a fixed inventory of possible speech sounds? The answer is no” (2005: 927). They go on to observe that “[T]he IPA makes no claims about the limits of the phonetic space nor does it posit any fixed number of possible phonetic distinctions” (Port and Leary 2005: 927). For example, Ladefoged and Maddieson (1996: 2–6) do not claim that it is possible to describe a closed set of “phonetic capabilities” of the human species, but hope that their continuous acoustic and articulatory parameters will be sufficient to differentiate all of those that appear. This points out the fundamental question posed in the present study: Are nasalized fricatives a phonetic capability of the human species? The conclusion is that they are, with a number of aerodynamic, acoustic, and potentially perceptual caveats. Nasalized fricatives, whether phonemic or potentially phonemic, are found at the edges of the expanding universe of the IPA.

Port and Leary further opine:

¹¹In practice, of course, detail in IPA transcriptions varies depending on the purpose of the transcription rather than some objective standard on how closely it should match the acoustic or articulatory reality of the utterance.

Back in the 1960s, it might have been reasonable to hope that phonetics research would gradually converge toward a fixed universal inventory of features, a limited set of vowel types, for example, that would be combinable into all words in all languages. But it is clear instead that forty years of phonetics research have provided absolutely no suggestion of convergence on a small universal inventory of phonetic types. Quite the opposite: the more research we do, the more phonetic differences are revealed between languages. So the hypothesis of a universal phonetic inventory should have been abandoned long ago on the basis of phonetic data (2005: 952).

They provocatively conclude: “There is no discrete universal phonetic inventory and thus phonology is not amenable to formal description” (2005: 953). While this statement is far too sweeping to accept at face value,¹² it points out the tension described earlier between the subphonemic and the phonemic in human language. It seems there is a necessity to distinguish between phonetic universals and phonemic universals. The present study, along with work by (Gerfen 1999, 2001), suggest that it is possible for nasal airflow and oral frication to occur simultaneously. The catch is that the spectral properties of the oral frication are so modified as to make the sound less distinct. While a ban on nasalized fricatives is not a phonetic universal, it seems like a plausible phonemic one, at least based on the grammatical sketches of languages in which they are claimed to occur (see Section 1.7).

So, does anything constrain the IPA from expanding, i.e. is the set of all linguistic sounds truly infinite? While Ladefoged (1990: 69) surmises that “[a] very substantial proportion of the possible sounds of the world’s languages have now been recorded” this does not imply that all the phonetic universals have been hammered out. Clearly, the matter of nasalized fricatives has been only partly resolved here. Lindblom (1990) takes the view that any explanation as to why possible speech sounds are or are not used in actual languages should come from outside linguistics. As Ladefoged (1990: 70) summarized well, “An explanation of something is an account of that event in terms of general principles that are not themselves dependent on the event” (Ladefoged 1990: 70). For nasalized fricatives, the reason for their subphonemic status likely has to do with the altered acoustics based on nasalization. Nevertheless, as discussed earlier, /š/ and /s/ could be phonemic in languages that lack non-sibilant fricatives. Such a phonemic distinction does not happen to occur in any known language, however. With this perplexity in mind, I conclude with Ladefoged:

¹²Formal descriptions can in fact include gradient dimensions, a possibility that Port and Leary (2005) unfortunately do not contemplate.

“We are at the moment a long way from being able to show whether the set of possible speech sounds is finite or not, and whether it has a particular form” (1990: 70).

The only true limits of the expanding IPA are the laws of physics (especially fluid dynamics and acoustics), the morphology of the human vocal tract, and constraints on the human auditory system (including the central nervous system that relays messages from the ear to the brain). Everything else is debatable.

4.6.2 The IPA as a Cartesian coordinate system

Because of its application to the IPA, it may be helpful to review the concept of the Cartesian product (Taylor 1999). The Cartesian product of two sets X and Y (also called the product set, set direct product, or cross product) is defined to be the set of all points (x, y) where $x \in X$ and $y \in Y$. It is denoted $X \times Y$. Expressed formally

$$X \times Y = \{(x, y) | x \in X \text{ and } y \in Y\} \quad (4.1)$$

This is called the Cartesian product since it originated in Descartes’ formulation of analytic geometry. In the Cartesian view, points in the plane are specified by their vertical and horizontal coordinates, with points on a line being specified by just one coordinate.

A quick glance at the consonant chart of the IPA may lead the casual observer to believe it is a kind of vectorized matrix where each phonetic symbol is defined as the Cartesian product $P \times M$ where $P = Place$ and $M = Manner$.¹³ However, every possible outcome of the equation is not listed in the chart. By convention, empty boxes indicate possible sounds that have not been observed. Shaded boxes indicate an impossible outcome.

The impossibility of a certain product $P_i \times M_j$ is determined based on the incompatibility of $Place_i$ and $Manner_j$, e.g. VELAR and TRILL. It is important to note that the basis of this judged incompatibility is in some cases physiological (velar trills) and aerodynamic (voiced pharyngeal plosives) but never acoustic or perceptual. It is perhaps the case that our grasp of the vocal tract’s morphology (with the application of a few basic aerodynamic principles) is more complete than our grasp of its acoustics. Last of all, our understanding of perception is still, I believe, in its early stages.

Thus, the (relative) morphological invariance of the human vocal tract should be

¹³For the consonant chart, X and Y are categorical variables, whereas for the vowel chart, they are continuous variables where $X = F_1$ and $Y = F_2$ or perhaps $F_2 - F_1$. For present purposes, the discussion will be limited to the product $P \times M$, though it has application to the product $F1 \times F2$, as well.

(and traditionally has been) a good starting point for discussions of phonetic and phonological universals. For precisely this reason, the standard division of consonants is by *Place* and *Manner*.

However, not all known phonemic possibilities can be reached using this product. For example, it is well known that some consonants have double articulations or the product $Place \times Place$. All the physiological possibilities for place of double-articulations can be investigated and then multiplied by manner (e.g. there are doubly articulated stops like $[\widehat{kp}]$ as well as fricatives like the simultaneously post-alveolar and velar $[\text{ɟʝ}]$).

Additionally, a few sounds may be said to have two manners of articulation, i.e. $Manner \times Manner \times Place$. The lateral fricatives $[\text{ɬ } \text{ɮ}]$ are two examples of $Manner \times Manner$ that happen to occur at the alveolar place of articulation. More germane to the present topic is the combination of manners NASAL and FRICATIVE, e.g. $[\text{ẽ } \text{f̃ } \text{θ̃ } \text{x̃}]$.

4.6.3 Nasalized fricatives: Shaded or empty cell?

One of the duties of the laboratory phonologist or experimental phonetician is to explain why some of the cells in the IPA chart are blank. In other words, why do some sounds that are judged to be physiologically possible fail to phonologize in any language of the world? While some of these omissions may be random, based on evolutionary luck of the draw, often the reasons are based on acoustic and perceptual principles. For example, what would a pharyngeal tap sound like? Could it be perceived in contrast to taps at other locations?

The problem of nasalized fricatives may be distilled to the following: should its cell in the IPA chart¹⁴ be shaded or empty? The results of the present study suggest it should be empty. Is this based on mere chance or on reduced perceptual salience? As I have discussed, the reason appears to depend on the fricative inventory of the language and on the fricatives that are singled out for nasalization.

Reports of nasalized fricatives cannot establish the sound as anything more than ‘rare’ in the vocal repertoire of the human species. Still, its existence highlights the importance of considering even the rarest of possibilities in determining phonetic universals. As Ladefoged and Everett have observed,

¹⁴This ‘cell’ is unfortunately a hypothetical one, since the arrangement of the consonant chart addresses $Manner \times Manner$, only in an *ad hoc* way, as for the lateral fricatives which are regarded as a single MANNER.

[W]e can never really tell what features will be needed for describing languages. In principle it is the complete set of human vocal sounds that can be integrated into the flow of speech, and that are sufficiently distinct from one another; but this is too cumbersome a notion to be of practical value for working linguists describing languages” (Ladefoged and Everett 1996: 799).

According to Ladefoged and Everett (1996), ‘central’ sounds are widely observed among the world’s languages and participate in many phonological processes, while ‘peripheral’ sounds are just the opposite. The authors meditate on the question of whether a universal feature set needs to be sufficiently powerful to account for phonetic rarities. They conclude that,

“Only through the close investigation of endangered and less well known languages will we be able to gather data that will help distinguish the two types of features, those required for widespread phonological processes, and those that specify phonetic rarities” (Ladefoged and Everett 1996: 799–800).

The results of the present study highlight this fact: by pursuing lines of inquiry to their logical conclusion, using instrumental means, we may come to learn new and surprising details about the development and phonologization of sounds, such as nasalized fricatives. In this regard, I agree with Port and Leary, who persuade their readers, “In a linguistics committed to the physical world (rather than to some Platonic heaven), language needs to be naturalized so as to fit into a human body. That implies, first of all, casting it into the realm of space and time” (Port and Leary 2005: 956). While aerodynamic principles suggest that nasalized fricatives cannot occur, this ultimately depends on one’s definition of ‘fricative’, which has to do with the acoustic nature of a sound and its phonological behavior. Phonetics is a science of gradient entities: individual phones naturally blend at the edges. Phonology, too, may not be discrete and ‘atomic’, as Port and Leary argue. Fricatives may be characterized by a range of gradient spectral properties and still, in the estimation of some, be considered fricatives. While voiceless nasalized fricatives appear to suffer the acoustic and potentially perceptual costs of nasalization (also a gradient phenomenon), it does not appear that they cease to be fricatives.

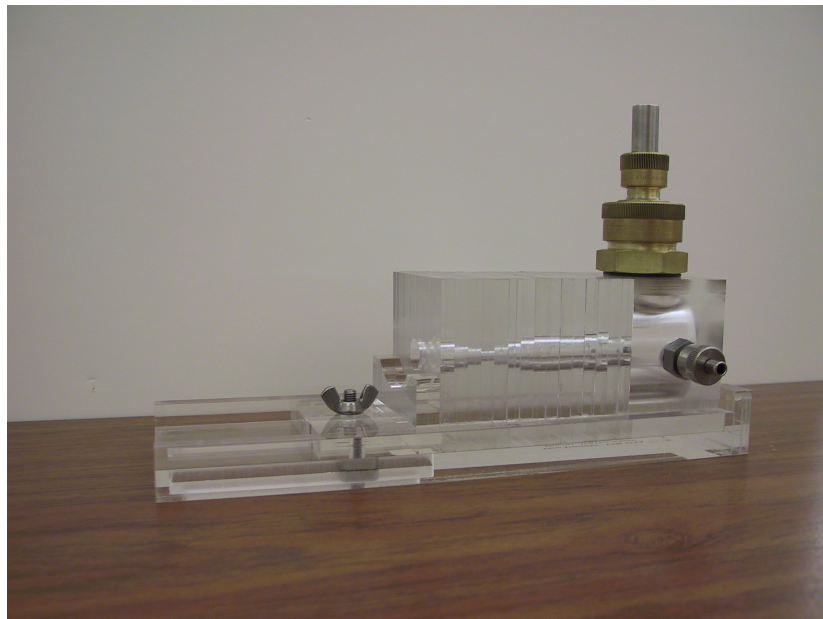


Figure 4.1: Photograph of the mechanical fricative model. The visible constrictions in the fricative model those of an American English alveolar [s] (Narayanan et al. 1995). The brass vent on top connected to the tube that served as the pseudo-velopharyngeal port. The metal tube at the side is for the measurement of pressure using a digital manometer. On the opposite side (not visible) there is a similar metal tube that may be attached to an air supply.

Bibliography

- Ali, L., R. Daniloff, and R. Hammarberg (1979). Intrusive stops in nasal-fricative clusters: An aerodynamic and acoustic investigation. *Phonetica* 36, 85–97.
- Anderson, S. R. (1975). The description of nasal consonants and internal structure of segments. In C. A. Ferguson, L. M. Hyman, and J. J. Ohala (Eds.), *Nasálfest: Papers from a Symposium on Nasals and Nasalization*, pp. 1–26. Stanford, CA: Language Universals Project.
- Badin, P. (1989). Acoustics of voiceless fricatives: Production theory and data. *Speech Transmission Laboratory Quarterly Progress and Status Report* 3, 33–55.
- Beasley, D. and K. Pike (1957). Notes on Huambisa phonemics. *Lingua Posnaniensis* 6, 1–8.
- Beddor, P. S. (1983). *Phonological and phonetic effects of nasalization on vowel height*. Ph. D. thesis, University of Minnesota. Bloomington, IN: Indiana University Linguistics Club.
- Beddor, P. S. (1993). The perception of nasal vowels. In M. K. Huffman and R. A. Krakow (Eds.), *Nasals, Nasalization, and the Velum*, Volume 5 of *Phonetics and Phonology*, pp. 171–196. San Diego: Academic Press.
- Bell-Berti, F. (1980). A spatial-temporal model of velopharyngeal function. In N. J. Lass (Ed.), *Speech and language: Advances in basic research practice*, pp. 291–316. New York: Academic Press.
- Bell-Berti, F. (1993). Understanding velic motor control: Studies of segmental context. In M. K. Huffman and R. A. Krakow (Eds.), *Nasals, Nasalization, and the Velum*, Volume 5 of *Phonetics and Phonology*, pp. 63–86. San Diego: Academic Press.

- Bell-Berti, F. and T. Baer (1983). Velar position, port size, and vowel spectra. *Proceedings of the 11th International Congress of Acoustics* 4, 19–21.
- Bell-Berti, F., T. Baer, K. S. Harris, and S. Niimi (1979). Coarticulatory effects of vowel quality on velar function. *Phonetica* 36, 187–193.
- Bell-Berti, F. and R. A. Krakow (1991). Anticipatory velar lowering: A coproduction account. *Journal of the Acoustical Society of America* 90, 112–123.
- Bergsveinsson, S. (1941). *Grundfragen der isländischen Satzphonetik*. Copenhagen: Metten & Co.
- Bhat, D. N. S. (1975). Two studies on nasalization. In C. A. Ferguson, L. M. Hyman, and J. J. Ohala (Eds.), *Nasálfest: Papers from a Symposium on Nasals and Nasalization*, pp. 333–352. Stanford, CA: Language Universals Project.
- Bladon, A., C. Clark, and K. Mickey (1987). Production and perception of sibilant fricatives: Shona data. *Journal of the International Phonetic Association* 17, 39–65.
- Bloch, B. (1950). Studies in colloquial Japanese, IV: Phonemics. *Language* 26, 86–125.
- Bognar, E. and H. Fujisaki (1986). Analysis, synthesis, and perception of the French nasal vowels. In *Proceedings of the International Conference on Acoustics, Speech, and Signal Processing*, Tokyo, pp. 1601–1604.
- Browman, C. P. and L. Goldstein (1986). Towards an articulatory phonology. *Phonology Yearbook* 3, 219–252.
- Brown, M. A., M. B. Jacobs, and R. Pelayo (1995). Adult obstructive sleep apnea with secondary enuresis. *Western Journal of Medicine* 163(5), 478–480.
- Brücke, E. (1856). *Grundzüge der Physiologie und Systematik der Sprachlaute für Linguisten und Taubstummenlehrer*. Vienna: Gerold.
- Brueckner, S. (2002). Crossing. Retrieved 07/14/05 from <http://www.mathworks.com/matlabcentral/files/2432/crossing.m>. Matlab Script.
- Carnochan, J. (1948). A study of the phonology of an Igbo speaker. *Bulletin of the School of Oriental and African Studies* 22, 416–427.

- Catford, J. C. (1977). *Fundamental Problems in Phonetics*. Bloomington, IN: Indiana University Press.
- Chamora, B. and R. Hetzron (2000). *Inor*. Number 118 in Languages of the World / Materials. Munich: Lincom Europa.
- Chen, M. Y. (1995). Acoustic parameters of nasalized vowels in hearing-impaired and normal-hearing speakers. *Journal of the Acoustical Society of America* 98, 2443–2453.
- Chiba, T. and M. Kajiyama (1941). *The Vowel: Its Nature and Structure*. Tokyo: Kaiseikan.
- Chomsky, N. and M. Halle (1968). *The Sound Pattern of English*. Cambridge, MA: MIT Press.
- Cohn, A. C. (1993). The status of nasalized continuants. In M. Huffman and R. Krakow (Eds.), *Nasals, Nasalization, and the Velum*, Volume 5 of *Phonetics and Phonology*, pp. 329–367. San Diego: Academic Press.
- Coltman, J. W. (1968). Sounding mechanism of the flute and organ pipe. *Journal of the Acoustical Society of America* 44(4), 983–992.
- Cotes, J. E., D. J. Chinn, and M. R. Miller (2006). *Lung Function: Physiology, Measurement and Application in Medicine* (6 ed.). Oxford: Blackwell.
- Crothers, J. (1978). Typology and universals of vowel systems. In *Universals of Human Language*, Volume 2, Phonology, pp. 93–152. Stanford, CA: Stanford University Press.
- Czermak, J. N. (1869). *Wesen und Bildung der Stimm- und Sprachlaute*, pp. 76–104. Leipzig: Wilhelm Engelman. [1879].
- Dang, J., H. Seikacho, and K. Honda (1995). Local and global effects of the pyriform fossa on speech spectra. *Journal of the Acoustical Society of America* 98(5), 2931.
- Daugherty, R. L. and J. B. Franzini (1965). *Fluid Mechanics with Engineering Applications* (6 ed.). New York: McGraw-Hill.
- Delattre, P. (1954). Les attributs acoustiques de la nasalité vocalique et consonantique. *Studia Linguistica* 8, 103–109.

- Doebelin, E. O. (1983). *Measurement Systems: Application and Design* (3 ed.). London: McGraw-Hill International Book Co.
- Durie, M. (1985). *A Grammar of Acehnese*, Volume 112 of *Verhandeligen van het Koninklijk Instituut voor Taal-, Land- en Volkenkunde*. Dordrecht: Foris.
- Einarsson, S. (1940). Nasal + spirant or liquid in Icelandic. *The Journal of English and Germanic Philology* 34, 462–464.
- Fant, G. (1970). *Acoustic theory of speech production* (2nd ed.). The Hague: Mouton.
- Fougeron, C. and C. L. Smith (1999). French. In *Handbook of the International Phonetic Association*, pp. 78–81. Cambridge: Cambridge University Press.
- Fourakis, M. and R. F. Port (1986). Stop epenthesis in English. *Journal of Phonetics* 14, 197–221.
- Fritzell, B. (1969). The velopharyngeal muscles in speech: an electromyographic cinradiographic study. *Acta Otolaryngologica Supplement* 250, 1–81.
- Fujimura, O. (1962). Analysis of nasal consonants. *Journal of the Acoustical Society of America* 34, 1865–1875.
- Fujimura, O. and J. Lindqvist (1971). Sweep-tone measurements of vocal-tract characteristics. *Journal of the Acoustical Society of America* 49, 541–558.
- Gander, W. and W. Gautschi (2000). Adaptive quadrature revisited. *BIT* 40, 84–101.
- Gerfen, C. (1999). *Phonology and Phonetics in Coatzospan Mixtec*, Volume 48 of *Studies in Natural Language and Linguistic Theory*. Dordrecht: Kluwer.
- Gerfen, C. (2001). Nasalized fricatives in Coatzospan Mixtec. *International Journal of American Linguistics* 67, 449–466.
- Gibson, C. H. (1999). Introduction to turbulent flow and mixing. In J. A. Schetz and A. E. Fuhs (Eds.), *Fundamentals of Fluid Mechanics*, pp. 83–88. New York: John Wiley and Sons.
- Goldstein, M. E. (1976). *Aeroacoustics*. New York: McGraw-Hill.

- Gordon, R. G. (2005). *Ethnologue: Languages of the World* (15th ed.). Dallas, TX: SIL International. Online version: <http://www.ethnologue.com/>.
- Green, M. M. and G. E. Igwe (1963). *A Descriptive Grammar of Igbo*. Berlin: Akademie Verlag for Oxford University Press.
- Greenlee, M. and J. J. Ohala (1980). Phonetically motivated parallels between child phonology and historical sound change. *Language Sciences* 2(2), 283–308.
- Hajek, J. (1997). *Universals of Sound Change in Nasalization*. Oxford: Blackwell.
- Harms, P. L. (1985). Epena Pedee (Saija) nasalization. In R. Brend (Ed.), *From Phonology to Discourse: Studies in Six Colombian Languages*, pp. 13–18. Dallas: Summer Institute of Linguistics.
- Harms, P. L. (1994). *Epena Pedee Syntax*, Volume 4 of *Studies in the Languages of Colombia*. Arlington, TX: Summer Institute of Linguistics.
- Hattori, S., K. Yamamoto, and O. Fujimura (1958). Nasalization of vowels in relation to nasals. *Journal of the Acoustical Society of America* 30, 267–274.
- Hawkins, S. and K. N. Stevens (1985). Acoustic and perceptual correlates of the non nasal-nasal distinction for vowels. *Journal of the Acoustical Society of America* 77, 1560–1575.
- Henderson, J. B. (1984). *Velopharyngeal function in oral and nasal vowels: a cross-language study*. Ph. D. thesis, University of Connecticut, Storrs.
- Hetzron, R. and H. M. Marcos (1966). Des traits pertinents superposés en ennemor. *Journal of Ethiopian Studies* 4, 17–30.
- Highland Council (2004). 2001 census profile for Applecross. Retrieved 01/13/06 from <http://www.highland.gov.uk/plintra/iandr/cen/sz/applecross.htm>.
- Hixon, T. J., F. D. Minifie, and C. A. Tait (1967). Correlates of turbulence noise production for speech. *Journal of Speech and Hearing Research* 10, 133–140.
- Hoaglin, P. F. and D. C. Hoaglin (1981). *Applications, Basics, and Computing of Exploratory Data Analysis*. Boston: Duxbury.

- Holder, W. (1669). *Elements of Speech: An Essay of Inquiry into the Natural Production of Letters*. London: J. Martyn.
- Hollander, M. and D. A. Wolfe (1973). *Nonparametric Statistical Methods*. New York: Wiley.
- House, A. S. and K. N. Stevens (1956). Analog studies of the nasalization of vowels. *Journal of Speech and Hearing Disorders* 21, 218–232.
- Jesus, L. M. T. and C. H. Shadle (2002). A parametric study of the spectral characteristics of European Portuguese fricatives. *Journal of Phonetics* 30, 437–464.
- Johnson, K. (1997). *Acoustic and Auditory Phonetics*. Oxford: Blackwell.
- Krakow, R. A. (1993). Nonsegmental influences on velum movement patterns: Syllables, sentences, stress, and speaking rate. In M. K. Huffman and R. A. Krakow (Eds.), *Nasals, Nasalization, and the Velum*, Volume 5 of *Phonetics and Phonology*, pp. 87–118. San Diego: Academic Press.
- Kruskal, W. H. and W. A. Wallis (1952). Use of ranks in one-criterion variance analysis. *Journal of the American Statistical Association* 47, 583–621.
- Kurowski, K. M. and S. E. Blumstein (1993). Acoustic properties for the perception of nasal consonants. In M. K. Huffman and R. A. Krakow (Eds.), *Nasals, Nasalization, and the Velum*, Volume 5 of *Phonetics and Phonology*, pp. 197–224. San Diego: Academic Press.
- Ladefoged, P. (1971). *Preliminaries to Linguistic Phonetics*. Chicago: University of Chicago. Midway reprint 1981.
- Ladefoged, P. (1990). Some reflections on the IPA. *UCLA Working Papers in Phonetics* 74, 61–76.
- Ladefoged, P. (2006). Representing linguistic phonetic structure. Manuscript in progress at the time of his death. Retrieved 03/2006 from [www.linguistics.ucla.edu/people /ladefoged/PhoneticStructure.pdf](http://www.linguistics.ucla.edu/people/ladefoged/PhoneticStructure.pdf).
- Ladefoged, P. and D. Everett (1996). The status of phonetic rarities. *Language* 72, 794–800.
- Ladefoged, P. and I. Maddieson (1996). *The Sounds of the World's Languages*. Oxford: Blackwell.

- Lastra, Y. (1984). Chichimeco Jonaz. In M. S. Edmonson (Ed.), *Supplement to the Handbook of Middle American Indians*, Volume 2, pp. 20–42. Austin, TX: University of Texas.
- Lilliefors, H. (1967). On the kolmogorov-smirnov test for normality with mean and variance unknown. *Journal of the American Statistical Association* 62, 399–402.
- Lindblom, B. (1990). On the notion of ‘possible speech sound’. *Journal of Phonetics* 18, 135–152.
- Liss, J. M. (1990). Muscle spindles in the human levator veli palatini and palatoglossus muscles. *Journal of Speech and Hearing Research* 33, 736–746.
- Lubker, J. F. (1968). An electromyographic-cinéradiographic investigation of velar function during normal speech production. *Cleft Palate Journal* 5, 1–18.
- Lubker, J. F., J. Lindqvist, and B. Fritzell (1972). Some temporal characteristics of velopharyngeal muscle function. In *Phonetics Symposium*. University of Essex Language Center.
- Lubker, J. F. and K. May (1973). Palatoglossus function in normal speech production. In *Papers from the Institute of Linguistics*, Volume 17, pp. 17–26. University of Essex.
- MacMahon, M. K. C. (1996). Phonetic notation. In P. T. Daniels and W. Bright (Eds.), *The World’s Writing Systems*, pp. 821–846. New York: Oxford University Press.
- Maddieson, I. (1983). The analysis of complex phonetic elements in Bura and the syllable. *Studies in African Linguistics* 14, 285–310.
- Maddieson, I. (1997). Phonetic universals. In W. Hardcastle and J. Laver (Eds.), *The Handbook of Phonetic Sciences*. Oxford: Blackwell.
- Maeda, S. (1993). Acoustics of vowel nasalization and articulatory shifts in French nasal vowels. In M. K. Huffman and R. A. Krakow (Eds.), *Nasals, Nasalization, and the Velum*, Volume 5 of *Phonetics and Phonology*, pp. 147–170. San Diego: Academic Press.
- Miller, G. A. and P. Nicely (1955). An analysis of perceptual confusions among some English consonants. *Journal of the Acoustical Society of America* 27(2), 338–352.
- Moll, K. L. (1960). Cinefluorographic techniques in speech research. *Journal of Speech and Hearing Research* 3, 227–241.

- Moll, K. L. (1962). Velopharyngeal closure on vowels. *Journal of Speech and Hearing Research* 17, 30–77.
- Moll, K. L. and R. G. Daniloff (1971). Investigation of the timing of velar movements during speech. *Journal of the Acoustical Society of America* 50, 678–684.
- Moll, K. L. and T. H. Shriner (1967). Preliminary investigation of a new concept of velar activity during speech. *Cleft Palate Journal* 4, 58–69.
- Narayanan, S., A. Alwan, and K. Haker (1995). An articulatory study of fricative consonants using MRI. *Journal of the Acoustical Society of America* 98(3), 1325–1347.
- Nusbaum, E. A., L. Foley, C. Wells, and L. S. Judson (1935). Experimental studies of the firmness of the velar-pharyngeal occlusion during the production of the English vowels [u i o e a ɔ æ]. *Speech Monographs* 2, 71–80.
- Ohala, J. J. (1975). Phonetic explanations for nasal sound patterns. In C. A. Ferguson, L. M. Hyman, and J. J. Ohala (Eds.), *Nasálfest: Papers from a Symposium on Nasals and Nasalization*, pp. 289–316. Stanford, CA: Language Universals Project.
- Ohala, J. J. (1983). The origin of sound patterns in vocal tract constraints. In P. F. MacNeilage (Ed.), *The Production of Speech*, pp. 189–216. New York: Springer-Verlag.
- Ohala, J. J. (1993). Sound change as nature’s speech perception experiment. *Speech Communication* 13, 155–161.
- Ohala, J. J. (1995). A probable case of clicks influencing the sound patterns of some european languages. *Phonetica* 52, 160–170.
- Ohala, J. J. (1996). Speech perception is hearing sounds, not tongues. *Journal of the Acoustical Society of America* 99, 1718–1725.
- Ohala, J. J. and M. Amadaor (1981). Spontaneous nasalization. *Journal of the Acoustical Society of America* 68, S54–S55. Abstract.
- Ohala, J. J. and M. Ohala (1993). The phonetics of nasal phonology: Theorems and data. In M. K. Huffman and R. A. Krakow (Eds.), *Nasals, Nasalization, and the Velum*, Volume 5 of *Phonetics and Phonology*, pp. 225–249. San Diego: Academic Press.

- Ohala, J. J., M.-J. Solé, and G. Ying (1998). Do nasalized fricatives exist? *Journal of the Acoustical Society of America* 103(5), 3085.
- Ohala, M. (1991). Phonological areal features of some Indo-Aryan languages. *Language Science* 13, 107–124.
- Padgett, J. (1991). *Stricture in Feature Geometry*. Ph. D. thesis, University of Massachusetts, Amherst.
- Padgett, J. (1997). Perceptual distance of contrast: Vowel height and nasality. In R. Walker, M. Katayama, and D. Karvonen (Eds.), *Phonology at Santa Cruz*, Volume 5, pp. 63–78. Santa Cruz, CA: University of California, Santa Cruz.
- Pétursson, M. (1973). Phonologie des consonnes nasales en islandais modernes. *La Linguistique* 9, 115–138.
- Pike, K. (1948). *Tone Languages*. Ann Arbor, MI: University of Michigan Press.
- Poirot, J. (1924). Sur l’articulation des nasales islandaises. In *Mélanges offerts à M. Charles Andler par ses amis et ses élèves*, Publications de la Faculté des Lettres de l’Université de Strasbourg, pp. 285–290. Strasbourg: Oxford.
- Port, R. F. and A. Leary (2005). Against formal phonology. *Language* 72, 927–964.
- Rabiner, L. R. and R. W. Schafer (1978). *Digital Processing of Speech Signals*. Upper Saddle River, NJ: Prentice-Hall.
- Rothenberg, M. (1977). Measurement of airflow in speech. *Journal of Speech and Hearing Research* 20, 155–176.
- Ruhlen, M. (1975). Patterning of nasal vowels. In C. A. Ferguson, L. M. Hyman, and J. J. Ohala (Eds.), *Nasálfest: Papers from a Symposium on Nasals and Nasalization*, pp. 333–352. Stanford, CA: Language Universals Project.
- Ruhlen, M. (1978). Nasal vowels. In *Universals of Human Language*, Volume 2, Phonology, pp. 203–242. Stanford, CA: Stanford University Press.
- Saksena, B. R. (1971). *Evolution of Awadhi*. Delhi: Motilal Banarsidas.

- Schadeberg, T. C. (1982). Nasalization in Umbundu. *Journal of African Languages and Linguistics* 4, 109–132.
- Shadle, C. H. (1983). Experiments on the acoustics of whistling. *The Physics Teacher* 21, 148–154.
- Shadle, C. H. (1985). The acoustics of fricative consonants. RLE Technical Report 506, Massachusetts Institute of Technology, Cambridge, MA.
- Shadle, C. H. (1997). The aerodynamics of speech. In W. J. Hardcastle and J. Laver (Eds.), *The Handbook of Phonetic Sciences*, Chapter 2, pp. 33–64. Oxford: Blackwell.
- Shepard, R. N. (1972). Psychological representation of speech sounds. In E. E. David and P. B. Denes (Eds.), *Human Communication: a Unified View*, pp. 67–113. New York: McGraw-Hill.
- Shosted, R. K. (2006a). Vowel context as a condition for nasal coda emergence: Aerodynamic evidence. *Journal of the International Phonetic Association* 36, 39–58.
- Shosted, R. K. (2006b). Whistled fricatives [ʂ ʐ] in Bantu: Acoustic origins. Presented at the Annual Meeting of the Linguistics Society of America, Albuquerque, New Mexico, January 7.
- Shosted, R. K. and B. Willgoos (2006). Nasals unplugged: The aerodynamics of nasal de-occlusivization in Spanish. In M. Díaz-Campos (Ed.), *Selected Proceedings of the 2nd Conference on Laboratory Approaches to Spanish Phonetics and Phonology*, pp. 14–21. Somerville, MA: Cascadilla Proceedings Project.
- Solé, M.-J. (1999). The phonetic basis of phonological structure: The role of aerodynamic factors. In *Proceedings of the 1st Congress of Experimental Phonetics*, Tarragona, Spain, pp. 77–94.
- Stevens, K. N. (1998). *Acoustic Phonetics*. Cambridge, MA: MIT Press.
- Stringer, M. and J. Hotz (1973). Waffa phonemes. In H. McKaughan (Ed.), *The Languages of the Eastern Family of the East New Guinea Highland Stock*, pp. 523–529. Seattle: University of Washington.

- Sturtevant, E. H. (1940). *Pronunciation of Latin and Greek*. New Haven, CT: Yale University.
- Sundberg, J. (1972). An articulatory interpretation of the 'singing formant'. *Speech Transmission Laboratory / Quarterly Progress Status Report, Stockholm 1*, 45–53.
- Taylor, P. (1999). *Practical Foundations of Mathematics*. Cambridge University Press.
- Ternes, E. (1989). *The Phonemic Analysis of Scottish Gaelic: Based on the Dialect of Applecross, Ross-shire* (2nd ed.). Hamburg: Helmut Buske Verlag.
- Thurneysen, R. (1946). *A Grammar of Old Irish*. Dublin: Dublin Institute for Advanced Studies.
- Trigo, R. L. (1988). *On the Phonological Derivation and Behavior of Nasal Glides*. Ph. D. thesis, Massachusetts Institute of Technology.
- Vance, T. (1987). *An Introduction to Japanese Phonology*. Albany: SUNY Press.
- Walker, R. (2000). *Nasalization, Neutral Segments, and Opacity Effects*. New York: Garland.
- Warren, D. W., R. M. Dalston, and R. Mayo (1993). Aerodynamics of nasalization. In M. K. Huffman and R. A. Krakow (Eds.), *Nasals, Nasalization, and the Velum*, Volume 5 of *Phonetics and Phonology*, pp. 119–146. San Diego: Academic Press.
- Warren, D. W. and A. B. Dubois (1964). A pressure-flow technique for measuring velopharyngeal orifice area during continuous speech. *Cleft Palate Journal 1*, 52–71.
- Weinberg, B. and Y. Horii (1975). Acoustic features of pharyngeal /s/ fricatives produced by speakers with cleft palate. *Cleft Palate Journal 12*, 12–16.
- Williamson, K. (1969). Ìgbò. In E. Dunstan (Ed.), *Twelve Nigerian languages*, pp. 85–96. London: Longmans.
- Wright, J. T. (1986). The behavior of nasalized vowels in perceptual vowel space. In J. J. Ohala and J. J. Jaeger (Eds.), *Experimental Phonology*, pp. 45–67. New York: Academic Press.

- Yu, A. C. L. (1999). Aerodynamic constraints on sound change: The case of syllabic sibilants. In J. O. et al. (Ed.), *Proceedings of the XIVth International Congress of Phonetic Sciences*, Volume 1, pp. 341–344.

Index

- Aceh, 17, 26
 Aerodynamic model, *see* Mechanical fricatives
 Aerodynamic signals
 Calibration, 30–33, 57–59, 61
 Numerical integration, 73–74
 Polynomial fitting, 71–73
 Apinayé, 3, 47
 Applecross Scots Gaelic, *see* Scots Gaelic, Applecross
 Authors
 Brücke, Ernst, 17
 Czermak, Johann Nepomuk, 17
 Fujimura, Osama, 13, 14, 49
 Gerfen, Chip, 3, 24, 27–35, 37, 44, 46, 50, 79, 107, 110, 112
 Hotz, Joyce, 3, 24, 27, 41–42, 110
 Jesus, Luis M. T., 64–68, 70, 85, 92, 98, 101, 102, 105
 Johnson, Keith, vii, 6, 13, 14, 48, 102
 Ladefoged, Peter, 2, 27, 38, 39, 42, 106, 108, 109, 111–113, 115
 Leary, Adam, 110–112, 115
 Maddieson, Ian, vii, 2, 27, 38, 39, 42, 106, 108, 111
 Maeda, Shinji, 12, 23
 Ohala, John, vii, 15, 18–25, 27, 32–36, 41, 45, 110
 Ohala, Manjari, 18, 20–21, 24, 27, 33–36, 45, 51, 110
 Port, Robert F., 110–112, 115
 Schadeberg, Thilo C., 24–27, 34, 35, 41, 110
 Shadle, Christine H., 6, 10, 11, 48, 59, 64–68, 70, 85, 92, 98, 101, 102, 105
 Solé, Maria-Josep, 18, 21–23, 33–35, 37, 45, 48, 49
 Stevens, Kenneth N., 10–14, 17, 48, 49, 102
 Stringer, Mary, 3, 24, 27, 41–42, 110
 Ternes, Elmar, 20, 35–37, 110
 Walker, Rachel, 3, 14, 21, 22, 28, 29, 35, 41–46
 Yu, Alan C. L., 18, 22, 34, 35, 45
 Awadhi, 107
 Barasano
 Northern, 47
 Southern, 47
 Brazilian Portuguese, 18, 22, 50–53, 63, 82–84, 100
 Bribri, 47
 Burmese, 106, 107

- Cabécar
 Northern, 47, 105
 Southern, 47, 105
 Cayuvava, 47
 Chichimeco-Jonaz, 37
 Coatzospan Mixtec, *see* Mixtec, Coatzospan
 Cubeo, 47

 Desano, 47
 Dynamic amplitude, 67, 70, 96

 Ennemor, *see* Inor
 Epena Pedee, 38, 44
 Epera, 47, 105

 French, 18, 50–53, 67, 82, 83, 88

 Gbeya, 47, 105
 Gokana, 47, 105
 Greek, 107
 Guanano, 47
 Guaraní, 47, 105
 Guaymi, 47

 High frequency spectral slope, 68, 87, 93,
 94, 99
 Hindi, 18, 50–53, 67, 71, 74, 82, 83, 85, 88,
 89
 Hmong, 106

 Içuã Tupí, 47
 Iaai, 106
 Icelandic, 39–40, 44
 Ìgbò, 38–39, 47
 Inor, 40–41, 44
 Irish, 107

 Japanese, 41
 Jivaro (Shuar), 22

 Kaiwá, 47

 Languages
 Aceh, 17, 26
 Apinayé, 3, 47
 Awadhi, 107
 Barasano
 Northern, 47
 Southern, 47
 Brazilian Portuguese, 18, 22, 50–53,
 63, 82–84, 100
 Bribri, 47
 Burmese, 106, 107
 Cabécar
 Northern, 47, 105
 Southern, 47, 105
 Cayuvava, 47
 Chichimeco-Jonaz, 37
 Cubeo, 47
 Desano, 47
 Ennemor, *see* Inor
 Epena Pedee, 38, 44
 Epera, 47, 105
 French, 18, 50–53, 67, 82, 83, 88
 Gbeya, 47, 105
 Gokana, 47
 Greek, 107
 Guanano, 47
 Guaraní, 47, 105
 Guaymi, 47

- Hindi, 18, 50–53, 67, 71, 74, 82, 83, 85, 88, 89
- Hmong, 106
- Içuã Tupí, 47
- Iaai, 106
- Icelandic, 39–40, 44
- Ìgbò, 38–39, 47
- Inor, 40–41, 44
- Irish, 107
- Japanese, 41
- Jivaro (Shuar), 22
- Kaiwá, 47
- Mixtec
 - Atatlahuca, 47
 - Coatzospan, 27–33, 37–38, 47
 - Ocotepec, 47
- Oregon, 47
- Parintintin, 47
- Portuguese, Brazilian, *see* Brazilian Portuguese
- Scots Gaelic
 - Applecross, 3, 35–37, 44, 105
- Shiriana, 47
- Shona, 102
- Siriano, 47
- Sundanese, 2
- Tatuyo, 47
- Tshwa, 102
- Tucano, 47
- Tuyuca, 47
- Umbundu, 24–27, 41, 44
- Waffa, 3, 24, 27, 41–42
- Low frequency spectral slope, 68–69, 89, 95, 96
- Mechanical fricatives, 50, 59–61, 97–99
- Mixtec
 - Atatlahuca, 47
 - Coatzospan, 27–33, 37–38, 47
 - Ocotepec, 47
- Model fricatives, *see* Mechanical fricatives
- Nasal harmony, 1, 2, 46
 - and Coatzospan Mixtec, 28
 - and nasalized fricatives, 104, 105
 - in Applecross Scots Gaelic, 35
 - in Coatzospan Mixtec, 28, 29
 - in Inor, 41
 - in other languages, 3, 42–45
 - in Umbundu, 26
- Oregon, 47
- Parameterization, *see* Spectral parameterization
- Parintintin, 47
- Portuguese, Brazilian, *see* Brazilian Portuguese
- Scots Gaelic
 - Applecross, 3, 35–37, 44, 105
- Shiriana, 47
- Shona, 102
- Sibilants, 70, 104, 105, 108–109
- Siriano, 47
- Spectral averaging, 64–66
- Spectral parameterization, 66–70

Spectral peak bandwidth, 70, 79, 89–90,
97

Strident, 108

Sundanese, 2

Tatuyo, 47

Tshwa, 102

Tucano, 47

Tuyuca, 47

Umbundu, 24–27, 41, 44

Universals

Phonetic, 28, 109–115

Voiceless nasals, 106–108

Waffa, 3, 24, 27, 41–42

Zero-crossing rate, 63, 86–87, 93