

PIAROA VOICELESS STOPS AS PARTIAL UNDERGOERS OF NASAL HARMONY

Jorge Emilio Rosés Labrada¹, Matthew Faytak², Tyler Schnoor¹, Myriam Lapierre³, Lev Michael⁴

¹University of Alberta ²University at Buffalo ³University of Washington ⁴University of California, Berkeley
jrosesla@ualberta.ca, faytak@buffalo.edu, tschnoor@ualberta.ca, mylapier@uw.edu, levmichael@berkeley.edu

ABSTRACT

Phonological analyses of long-distance nasal harmony (LDNH) processes classify segments as triggers, undergoers, blockers, or transparent, based primarily on researchers' subjective judgments of nasality. Using novel field instrumental data from Piaroa (a Jod̥i-Sáliban language with LDNH), we investigate stop duration, voice onset time (VOT), and nasal and oral airflow during the Piaroa voiceless stops in all-oral environments and continuous nasal harmony spans. We find that nasal harmony has no effect on stop duration and VOT, but voiceless stops in nasal contexts exhibit nasal airflow above an oral baseline at the onset of closure, suggesting that voiceless stops in nasal spans are partial undergoers of LDNH. This partial undergoer behavior is consistent with a model of coactivation of antagonistic gestural specifications for velum activity of the voiceless stop undergoer and the nasal harmony span. Similar voiceless obstruent pre-nasalization is likely widespread cross-linguistically but vastly underreported due to its subtle acoustic effects.

Keywords: nasal harmony; voiceless stops; voice onset time; nasal airflow; Piaroa

1. INTRODUCTION

Piaroa (Glottocode: piar1243; ISO 639-3: pid) is a Jod̥i-Sáliban language spoken along the Middle Orinoco River and its tributaries in an area that straddles the border between Colombia and Venezuela. Like numerous other South American languages [1, p. 268], Piaroa exhibits long-distance nasal harmony (LDNH). In this paper, we investigate the effect of LDNH on Piaroa voiceless stops, specifically examining their overall duration and their voice onset time (VOT), as well as the rate of nasal and oral airflow during stop production.¹

1.1. Nasal harmony and consonants

In nasal harmony systems, segments can act either as *triggers*, spreading their nasality to other segments; *undergoers*, undergoing nasalization; *blockers*, stopping the spread of nasality; or *transparent*, neither impeding the spread of nasality nor being

(phonologically) affected by it. Which (classes of) segments will exhibit which behaviors is, to some degree, language-dependent. The behavior of voiceless obstruents in particular tends to vary between blocking and transparency. For example, voiceless stops are blockers in Warao [2], but have been characterized as (phonologically) transparent in Paraguayan Guaraní [3].

However, there is some evidence that /p t k/ are affected by nasal spreading in Paraguayan Guaraní and are not fully phonetically transparent. Walker [3] finds that the extent of carryover voicing in intervocalic voiceless stops can be affected: the VOT of /p/ and /t/, but not /k/, is significantly longer in nasal harmony spans. These effects could be caused by carryover nasal airflow into the stop closure as demonstrated for Guaraní in [4].

1.2. Piaroa consonants and Piaroa nasal harmony

The consonant inventory for Piaroa consists of 14 oral stops /p t k kʷ pʰ tʰ kʰ pʰ tʰ kʰ kʷʰ bʰ dʰ ʔ/, two affricates /tsʰ tʃʰ/, three fricatives /s h hʷ/, two nasal stops /m n/, two glides /w j/, and one rhotic /r/ [5], [6]. Krute [5, p. 62] describes a process of LDNH in Piaroa, which is triggered by phonemically nasal vowels in classifier suffixes and targets preceding vowels and voiced consonants (i.e., voiced stops, glides, and the tap). He argues that nasality spreads leftward to these undergoers, which are realized as their corresponding nasalized allophones, up to the first voiceless consonant preceding the trigger. That is, Piaroa voiceless consonants are *blockers* according to this account.

Recently collected primary fieldwork data, however, shows that Piaroa voiceless stops do not act as blockers, but rather appear to be phonologically transparent. In particular, in verb forms with a masculine classifier /-a/ or with the durative suffix /-æ:/ [7], voiceless stops do not impede the leftward spread of nasalization and they appear to be unaffected in their phonetic realization. To test the behavior of all segments (both consonants and vowels) in nasal harmony spans, instrumental phonetic data was collected in May 2022 with 10 Piaroa speakers from the community of Babel (Atures municipality, Amazonas state, Venezuela).

1.3. Research questions and hypothesis

In this paper, we focus on the behavior of Piaroa's plain voiceless stops /p t k k^w/, comparing their realization in nasal harmony spans with that outside such spans in an oral, non-nasal harmony environment, with an eye towards comparison with previous findings on Paraguayan Guaraní [3], [4]. We first consider the effects of nasality on overall voiceless stop duration and VOT duration. In addition, we examine nasal and oral airflow trajectories to evaluate the timing of the velum-raising gesture for voiceless stops and whether voiceless stops behave as transparent segments or undergoers of LDNH. Our research questions are the following:

1. Is there a difference in stop duration or VOT for voiceless stops in a nasal harmony span vs. in an oral environment?
2. During the production of plain voiceless stops in a nasal harmony span, is the level of nasal airflow above that within an oral, non-LDNH environment?

2. METHODOLOGY

The data reported here is from two speakers (1F, speaker JPP; 1M, speaker ROS) and was recorded in a quiet room in Puerto Carreño, Colombia.

2.1. Stimuli and data collection

A wordlist that included two target verbs containing each consonant in the inventory was created by the first author with help from Piaroa native speakers. The male future tense conjugation employs a male classifier suffix,² which contains a phonemic nasal vowel /-α/ and triggers leftward LDNH; while the female future tense form of the verb has the female classifier suffix /-æhu/, which contains only oral segments and does not trigger LDNH [7]. All target productions were thus future-tense verb forms, which allowed us to record stimulus pairs with the target segments occurring in similar segmental and prosodic contexts but differing in nasality, as illustrated by contrasting the forms in Table 1.

Oral and nasal airflow were collected using a modified Laryngograph® D-800 electroglottograph and a Glottal Enterprises oro-nasal mask. Stimuli were presented in a different random order for each speaker. Speakers were verbally prompted in Spanish by the first author to provide the Piaroa masculine or feminine present-tense of a given verb (to confirm that the intended root had been identified); they then produced five repetitions of the target future-tense form. This resulted in 10 tokens per verb: 5 in the masculine form (i.e. nasal context) and 5 in the feminine form (i.e. oral context). The 10 words

selected for this study (Table 1) yielded 420 tokens (210 per speaker, Table 2). The absence of word-medial /k/ in the data is due to an apparent lexical gap.

Target	Inflected FUT forms	Root gloss
p-, -k ^w -	f. pa-ʔd-æk ^w -æhu-sæ m. pa-ʔn-æk ^w -α-sæ V-1SG-FUT-CLF-1	'sing, pray'
p-, -k ^w -	f. pæ-ʔd-æk ^w -æhu-sæ m. pæ-ʔn-æk ^w -α-sæ V-1SG-FUT-CLF-1	'say'
-p-, -k ^w -	f. hæ-ʔd-ep-æk ^w -æhu-sæ m. hæ-ʔn-ēp-æk ^w -α-sæ V-1SG-SUFF-FUT-CLF-1	'ask for'
t-, -k ^w -	f. tē-ʔd-æʔd-æk ^w -æhu-sæ m. tē-ʔn-æʔn-æk ^w -α-sæ V-1SG-SUFF-FUT-CLF-1	'open (hand)'
-t-, -k ^w -	f. tʃi-teh-æk ^w -æhu-sæ m. tʃi-tēh-æk ^w -α-sæ 1SG-V-FUT-CLF-1	'shine (light)'
-t-, -k ^w -	f. tʃ-αʔdit-æk ^w -æhu-sæ m. tʃ-αʔnīt-æk ^w -α-sæ 1SG-V-FUT-CLF-1	'work'
k-, -k ^w -	f. ke-ʔd-æʔd-æk ^w -æhu-sæ m. kē-ʔn-æʔn-æk ^w -α-sæ V-1SG-SUFF-FUT-CLF-1	'finish'
k-, -p-, -k ^w -	f. kæ-ʔd-ep-æk ^w -æhu-sæ m. kæ-ʔn-ēp-æk ^w -α-sæ V-1SG-SUFF-FUT-CLF-1	'lift'
k ^w -, -k ^w -	f. k ^w æ-ʔd-æk ^w -æhu-sæ m. k ^w æ-ʔn-æk ^w -α-sæ V-1SG-FUT-CLF-1	'kill/hit'
k ^w -, -k ^w -	f. k ^w o-ʔd-æk ^w -æhu-sæ m. k ^w ō-ʔn-æk ^w -α-sæ V-1SG-SUFF-FUT-CLF-1	'swim'

Table 1: Target words and segments analyzed here (V = verb root, f. = female form, m. = male form). All tokens translate to "I will V".

Segment	Word-initial	Word-medial
/p/	40	40
/t/	25	35
/k/	40	0
/k ^w /	40	200

Table 2: Number of tokens of each segment by word position and place of articulation.

2.2. Data processing and analysis

Data were segmented manually in Praat [8]. Start of closure and end of release were located using the drop (or rise) in the amplitude of low-frequency audio signal components. In the case of word-initial stops, speaker JPP did not pause between repetitions of stimuli, leading to a clearly identifiable closure

portion. Speaker ROS did pause between stimulus repetitions. To create suitable time-series data for GAMM analysis, the start of closure from ROS's word-initial stops was estimated from oral or nasal airflow (at the cessation of airflow corresponding to exhalation or inhalation). The timing of stop release was determined based on burst energy in the Praat spectrogram or, if this was not sufficient, from the appropriately timed spike in the oral airflow channel. Fig. 1 illustrates the outlined annotation scheme.

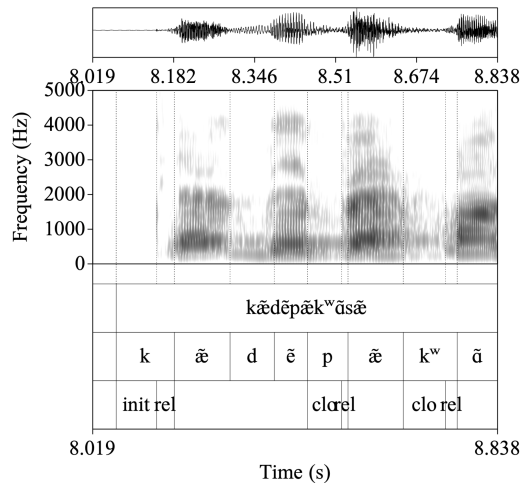


Figure 1: Annotated partial 'lift' male-form token by speaker ROS.

Using these reference points, VOT (release) and the duration of the entire stop (closure+release) were calculated using a Praat script. For 28 of the 420 stop tokens annotated, no clear release phase could be identified. These were excluded from VOT analysis but included in the stop duration analysis. Initial stops were excluded from analysis for stop duration. Speaker ROS's estimated initial closures were not used as a measure of duration and only used to define the time-series data used for GAMM models.

Stop duration and VOT data were submitted to separate linear mixed-effects regressions in R using *lme4* v1.1-31, with *p*-values calculated using *lmerTest* v3.1-3 [9], [10]. Two VOT models were constructed, one for word-medial stops and the other for word-initial stops. The stop duration model excluded all word-initial stops. Models assessed the dependent variable with respect to fixed effects of nasality, stop place, speaker,³ and their interactions, with random intercepts for word.

Oral and nasal airflow signals were extracted for all target stops using a custom Python script. Airflow signals were smoothed using a Butterworth filter implemented with the *butter* function in the Python *scipy* library (50 Hz cutoff, third-order), downsampled to 200 Hz, and submitted to generalized additive mixed models (GAMMs) built with *mgcv* in R [11]. GAMMs were chosen for their ability to visualize the entire trajectory of airflow over

the duration of the consonant, in line with other recent work on time-varying aerodynamic data [12], [13]. Separate GAMMs estimated oral and nasal flow over percent duration as a function of nasality, with factor smooths for speaker and position (initial vs. medial) and a random-effect smooth for word. Difference smooths were created using *tidymv* [14].

3. RESULTS

3.1. Stop duration and VOT

In reporting the results of the VOT and duration models, we focus on the presence or absence of main effects of nasality and its interactions with speaker and stop place. We begin with the (medial-only) stop duration model. Here, the main effect of nasality on segment duration fails to reach significance ($\beta = -0.00711$, $t = -1.054$, $p = 0.294$), along with all interactions of nasality with speaker and stop place. Turning to the word-medial VOT model, the main effect of nasality again fails to reach significance ($\beta = .00114$, $t = 0.303$, $p = 0.763$), along with all interactions of nasality with speaker and stop place.

In the word-initial VOT model, the main effect of nasality once again fails to reach significance ($\beta = 0.000495$, $t = -0.136$, $p = 0.894$). However, the *interaction* of nasality and Speaker ROS reaches significance ($\beta = -0.0165$, $t = -3.221$, $p = 0.00671$), suggesting a slight *shortening* effect of nasality on VOT which is particular to Speaker ROS in initial position. The three-way interaction of nasality, Speaker ROS, and /k^w/ is also weakly significant ($\beta = 0.0172$, $t = 2.372$, $p = 0.0339$), suggesting a moderating effect on the VOT of word-initial /k^w/ specifically in nasal spans.

To summarize, modeling suggests that Piaroa medial stop duration was not affected by nasal harmony spans, in line with Walker's [3] findings for Paraguayan Guaraní, where no effect of nasality on the duration of voiceless stops was observed. However, Piaroa medial stop VOT is likewise not affected, which is *not* in line with Walker's findings: VOT for Paraguayan Guaraní /p/ and /t/ (but not /k/) lengthened within a LDNH span.

3.2. Nasal and oral airflow trajectory

Since VOT and duration data offer indirect evidence for an impact of nasal harmony on voiceless stop production, we now turn to the airflow data for more direct evidence. GAMM fitted smooths for nasal and oral airflow are given in Fig. 2. In nasal harmony spans, compared to all-oral spans, Piaroa plain voiceless stops exhibit elevated nasal airflow from the onset of stop closure, persisting through the first third of the stop's duration. This continued airflow can be

understood as partial nasalization of the voiceless stop due to a delay in the stop's associated velum-raising gesture under the nasal harmony span. Demolin [4] observes that Paraguayan Guaraní voiceless stops in nasal spans exhibit a similar pattern, seemingly as a consequence of a spike in nasal airflow at the end of a nasalized vowel just before stop closure. While preceding vowels are beyond the scope of this paper, this suggests an avenue for further research.

Oral airflow has a much greater variance; in nasal harmony spans, it differs from that during oral, non-harmonized spans mainly at the onset of the stop closure, where fitted airflow in the oral span briefly exceeds fitted airflow in the nasal span. The source of this reduction in oral airflow is less obvious; we speculate that it may be a consequence of globally reduced oral airflow during nasal spans.

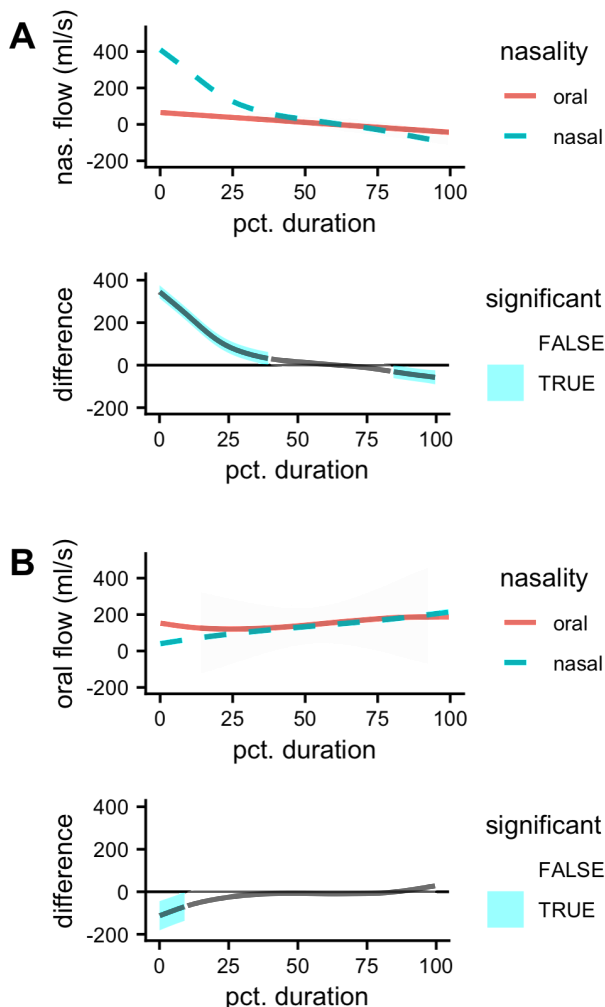


Figure 2: GAMM fits and difference smooths for nasal airflow (A) and oral airflow (B) in oral and nasal conditions.

4. DISCUSSION

While stop duration and VOT data did not suggest that nasal harmony spans affect the production of /p t k k^w/, our airflow data offer direct evidence that voiceless stops in Piaroa are partial undergoers of nasal harmony, with the initial portion of their closure realized with relatively high nasal airflow (in line with the Paraguayan Guaraní airflow data reported in [4]). This prenasalization process is observed in several other Amazonian indigenous languages. For instance, nasal airflow data for Panãra (Jê) shows that /p t s k/ are prenasalized immediately after phonemically nasal vowels [15]. Lapierre [15, p. 27] discusses additional cases of obstruent prenasalization and suggests that this phenomenon is especially common among Amazonian languages. We speculate that voiceless obstruent prenasalization is underreported in the literature due to its low perceptual salience.

That Piaroa voiceless stops are partial undergoers of LDNH has important implications for models of phonological representation, as it suggests that different portions of a segment (or subsegments) may not behave uniformly as triggers, undergoers, transparent segments or blockers within a LDNH system. In Piaroa, the initial portion of voiceless stop closure behaves as an undergoer, while the latter portion behaves as transparent. Similarly, complex nasal segments in Tupí-Guaraní languages may behave as partial triggers and partial blockers of LDNH [16]. These findings can be captured by formal phonological models of subsegmental representations, such as Q-Theory [15]–[17].

The Piaroa airflow data presented here are equally consistent with an account of partial gestural deactivation under conflicting gestural specifications from the voiceless stop and the nasal harmony span. Per Smith [18], when a segment specified with a velum-raising gesture occurs inside a nasal harmony span specified for a velum-lowering gesture, antagonism between the two gestures arises. Smith proposes that obstruent (partial) transparency (and partial undergoing) in nasal spreading results from favoring the obstruent's velum raising gesture. Specifically, the velum begins to raise during the production of oral constriction for the oral obstruent, and some nasal airflow persists into the stop closure. Upon stop release, the harmony span's velum lowering gesture is reactivated for any adjacent segments in the nasal span that are compatible with a lowered velum gesture. If this gestural account of obstruent transparency within LDNH is correct, we should indeed expect voiceless obstruents in many other languages to show a similar pattern of prenasalization as Piaroa.

5. REFERENCES

- [1] L. Campbell, "Typological characteristics of South American indigenous languages," in *The Indigenous languages of South America. A Comprehensive Guide*, L. Campbell and V. M. Grondona, Eds., Berlin; Boston, MA: De Gruyter Mouton, 2012, pp. 259–330.
- [2] L. Peng, "Nasal harmony in three South American languages," *International Journal of American Linguistics*, vol. 66, no. 1, pp. 76–97, 2000.
- [3] R. Walker, "Guaraní voiceless stops in oral versus nasal contexts: An acoustical study," *Journal of the International Phonetic Association*, vol. 29, no. 1, pp. 63–94, Jun. 1999.
- [4] D. Demolin, "The phonetics and phonology of nasal harmony and nasal segments in Guaraní," presented at the Nasal 2009, Montpellier, Jan. 2009. [Online]. Available: https://www.researchgate.net/publication/228985362_The_phonetics_and_phonology_of_nasal_harmony_and_nasal_segments_in_Guarani
- [5] L. D. Krute, "Piaroa nominal morphosemantics," PhD, Columbia University, 1989.
- [6] E. E. Mosonyi, "Elementos gramaticales del idioma piaroa," in *Lenguas indígenas de Colombia: una visión descriptiva*, M. S. González de Pérez and M. L. Rodríguez de Montes, Eds., Santafé de Bogotá: Instituto Caro y Cuervo, 2000, pp. 657–668.
- [7] J. E. Rosés Labrada, "The Sáliban languages," in *Amazonian Languages, An International Handbook*, P. Epps and L. Michael, Eds. Berlin and New York: De Gruyter Mouton, forthcoming.
- [8] P. W. Boersma, "Praat: doing phonetics by computer [Computer Program] Version 6.1.39." 2022.
- [9] D. Bates, M. Maechler, and B. Bolker, "lme4: Linear Mixed-Effects Models using 'Eigen' and S4. R package, version 1.1-31." 2022.
- [10] A. Kuznetsova, P. Brockhoff, R. Christensen, and S. Jensen, "lmerTest: Tests in Linear Mixed Effects Models. R package, version 3.1-3." 2020.
- [11] S. Wood, "Package 'mgcv'. R package, version 1.8-41." 2015.
- [12] H. M. Stoakes, J. M. Fletcher, and A. R. Butcher, "Nasal coarticulation in Bininj Kunwok: An aerodynamic analysis," *Journal of the International Phonetic Association*, vol. 50, no. 3, pp. 305–332, Dec. 2020, doi: 10.1017/S0025100318000282.
- [13] F. Desmeules-Trudel and M. Brunelle, "Phonotactic restrictions condition the realization of vowel nasality and nasal coarticulation: Duration and airflow measurements in Québécois French and Brazilian Portuguese," *Journal of Phonetics*, vol. 69, pp. 43–61, Jul. 2018, doi: 10.1016/j.wocn.2018.05.001.
- [14] S. Coretta, "Tidymv: tidy model visualisation for generalised additive models. R package, version 3.3.2." 2020.
- [15] M. Lapiere, "Towards a theory of subsegmental and subfeatural representations: The phonology and typology of nasality," PhD dissertation, University of California, Berkeley, Berkeley, CA, 2021. Accessed: Jan. 02, 2023. [Online]. Available: <https://escholarship.org/uc/item/9x90w7c9>
- [16] M. Lapiere and K. Russell, "Opposite directionality in [+/-F] agreement: The case of Tupí-Guaraní nasal harmony," presented at the 97th Annual Meeting of the Linguistic Society of America, Denver, CO, 2023.
- [17] S. S. Shih and S. Inkelas, "Autosegmental aims in Surface-Optimizing Phonology," *Linguistic Inquiry*, vol. 50, no. 1, pp. 137–196, Jan. 2018, doi: 10.1162/ling_a_00304.
- [18] C. Smith, "Nasal spreading as defective gestural deactivation," *Proceedings of the 14th Annual Meeting on Phonology*, vol. 2, no. 0, Art. no. 0, Jun. 2016, doi: 10.3765/amp.v2i0.3771.
- [19] J. E. Rosés Labrada, "The Piaroa subject marking system and its diachrony," *Journal of Historical Linguistics*, vol. 8, no. 1, pp. 31–58, 2018.

¹ This research was funded by the US National Science Foundation (award #1918064) and carried out under UC Berkeley IRB protocol 2019-06-12311 and University of Alberta REB protocol Pro00075931. We would like to thank the speakers of Piaroa from the community of Babel who participated in this research.

² For additional details about the future tense conjugation (including the use of the classifiers /-a/ and /-æhu/ as indexes of subject grammatical gender), see [19].

³ Note that speaker is coded as a fixed effect, rather than as a random effect, in all models owing to the small sample size.