

Nasal Coda Loss in the Chengdu Dialect of Mandarin: Evidence from RT-MRI

Sishi Liao¹, Phil Hoole¹, Conceição Cunha¹, Esther Kunay¹, Aletheia Cui¹, Lia Saki Bučar Shi-gemori¹, Felicitas Kleber¹, Dirk Voit², Jens Frahm², Jonathan Harrington¹

¹Institute of Phonetics and Speech Processing, University of Munich

²Max Planck Institute for Multidisciplinary Sciences

sishi.liao|hoole|cunha|lia|kleber|jmh@phonetik.uni-muenchen.de,
es.kunay|A.Cui@lmu.de, dvoit|j.frahm@mpinat.mpg.de

Abstract

In the Chengdu Dialect of Mandarin, the /*(V)an*/ rime words have been described to have undergone a nasal loss process in the last decades. However, no acoustical or physiological evidence has been provided so far. In this study, we investigate this sound change process by directly looking at the velum gesture in the *target segments* from 4 Chengdu speakers. By means of real-time Magnetic Resonance Imaging (rt-MRI), the velum opening signal was captured along with synchronized and noise suppressed audio. The maximum degree of velum opening was compared between tautosyllabic and heterosyllabic VN sequences for different vowels ($N = /n, \eta/$). Nasal consonant loss was most evident for tautosyllabic /*(V)an*/ rime words. This sound change, together with the observed diachronic vowel raising in /*(V)an*/ rimes, is compatible with other research showing a preference for low vowel raising before nasal consonants. This phonetically motivated oral vowel, which is a consequence of nasal coda loss and vowel raising, would form a new phonological contrast in this dialect e.g., from /*pa, pan*/ to /*pa, pɛ*/.

Index Terms: Nasal loss, rtMRI, Chengdu Dialect, velum gesture, sound change

1. Introduction

The Chengdu Dialect is a branch of southwestern Mandarin in China. It shares some phonological characteristics with Standard Mandarin (SM), yet it also has some unique phonetic qualities, one of which is the weakening of the nasal coda. It has been reported that the nasal coda in Chengdu Dialect is weakened especially after low vowels [1], [2]. An empirical investigation was undertaken to test whether there was any evidence of post-vocalic nasal consonant weakening and concomitant changes in the preceding vowel, in particular the degree of nasalization or changes in vowel height.

The production and perception of nasalization have been investigated in many other branches of Mandarin and some southern dialects in China through acoustic analysis [3], nasometric measurements [4], perception studies [5], and ultrasound imaging [6]. Except for ultrasound imaging, all methods use indirect indicators and measurements (formants of nasal and oral vowel, the energy ratio of nasal to nasal plus oral, human behavioral experiment results, etc.) to analyze the characteristics of nasal codas. In addition to ultrasound imaging, Magnetic Resonance Imaging (MRI) has been used for investigating speech production more directly since the last century [7]. More specifically, real-time MRI (rtMRI) has been applied in the study of movements of articulators during speech [8] and most recently to

further investigate the spatio-temporal patterns of velum kinematics [9], [10].

At least two pathways leading to the loss of the nasal consonant have been suggested. Pathway A starts from anticipatory coarticulation in the VN sequence by which the V is nasalized as $\tilde{V}$ and the whole sequence is phonetically realized as $\tilde{V}n$ before the nasal is deleted and only a nasalized vowel $\tilde{V}$ remains [11]. Pathway B also starts from anticipatory coarticulation in the VN sequence, but instead of leading to a contrastive nasalized vowel, the nasal is completely lost, and an oral vowel takes its place [12]. However, the nasal loss of VN in Chengdu Dialect only occurs in /*(V)an*/ rime words and is not motivated by the voicing status of a post-nasal consonant as in [11]–[14]. With pathway A, it has often been observed that vowels are simultaneously raised [15], [16]. Such vowel raising has also been described in the Chengdu Dialect [17]. However, whether the vowel raising co-occurs with nasal loss (pathway B) remains an open question.

Diachronic changes together with synchronic variations have made Chengdu /*(V)an*/ rimes quite different from other vowel-nasal combinations. The diachronic development of nasal coda in Chengdu Dialect was documented as early as 1900 [18] in which the /*(V)an*-rime words were still transcribed as /*an/-/ian/-/uan/-/yan*/ with a complete nasal coda. In the 1950s, the /*(V)an*-rime words by Chengdu native speakers born in 1930s were already transcribed as $[a^n]-[iɛ^n]-[ua^n]-[yɛ^n]$ and were documented accompanied by the observation: ‘the /*(V)an*/ rimes are now presented with a tongue gesture raising towards the alveolar ridge, instead of actually touching it’ [1], [19]. In the 1980s, the nasal codas were transcribed from speakers born in 1960s as $[\tilde{a}]-[\tilde{i}\tilde{e}]-[\tilde{u}\tilde{a}]-[\tilde{y}\tilde{e}]$ [2]. Finally in the 2010s, [20] suggested that these /*(V)an*-rime words were realized by some as plain oral vowels $[\tilde{a}]-[\tilde{i}\tilde{e}]-[\tilde{u}\tilde{a}]-[\tilde{y}\tilde{e}]$. Field work conducted at 150 sites in the Sichuan province, on the other hand, showed that, $[a^n]$, $[\tilde{a}]$ and $[\tilde{a}^n]$ were all possible realizations of /*(V)an*/ rime words [21].

The Chengdu Dialect has a phonological contrast between /*n*/ or /*η*/, both of which can occur in onset and coda positions. However, in this study, no words with onset /*η*/ are included due to its unstable production among individuals.

In this study, we explore the variation of maximum velum opening with signals derived from rt-MRI. The corpus consists of carrier phrases with *target segments* differing in vowel context, syllabification, place of articulation (POA) of the nasal, and vowel type. We aim to assess the extent to which vowel nasalization, vowel raising, and weakening of the nasal consonant are characteristics of Chengdu nasal rhymes. We aim to answer the question: to what extent do vowel context, POA,

vowel type and syllabification affect the nasality within a VN sequence in the Chengdu Dialect?

2. Methods

2.1. Participants and procedures

Four self-reported native Chengdu speakers (2 female) were recruited for this experiment. All participants were born in Chengdu city, Sichuan, China, and native speakers of the Chengdu Dialect. The age range of the participants was 22 to 33 years. They consented to the recording by signing the corresponding consent form prior to the experiment and they received compensations for their participation. Experimental procedures were approved by the ethics committee of the Faculty of Medicine at LMU Munich. All the prompts were sent via email to the participants together with online tutorials of the experimental setup for familiarization purpose.

Real-time MRI data were collected with a 3T MRI system at the Max Planck Institute for Multidisciplinary Sciences in Göttingen, Germany with the participant being in a supine position during the recordings and the head being placed in a head coil. The randomized prompts were split into sessions and were shown on a monitor, which was reflected through mirrors on the head coil to the eyes of the participants. Upon presentation of a prompt, participants were asked to read out loud the prompt at a normal pace. There was a little break between sessions if it was required by the participants.

2.2. Speech materials

In this corpus, the target characters were either CV₁ or CV₁N, varied in the onsets (C), rimes (V₁ or V₁N), and tones. The C included /p, p^h, f/, and there were also V₁ initial characters. V₁ was one of the simple /a, e, i/ or compound /ia, ua/ vowels. The nasal N differed in place of articulation (POA) and was either velar or alveolar. The tones included T1, T2, T3, T4. In the Chengdu Dialect, T1 is a rising tone, T2 and T3 are both falling tones starting from mid and high, and T4 is a falling-rising tone. Thus, the target characters were composed of one of the onsets /p, p^h, f, Ø/, the rimes in Table 1 (here represented in Pinyin, a romanization system of Mandarin), and one of the tones T1 – T4.

Table 1: *Rimes collected within the corpus with nasals in coda (tautosyllabic) and onset (heterosyllabic, following the rime) position; boxed rimes are analyzed in model 1, and grayed rimes in model 2, see section 3.2.*

Target Segment Type (Table 2)	/CV ₁ #nV ₂ /	/CV ₁ n#V ₂ /	/CV ₁ n#V ₂ /
SYLLABIFICATION	Hetero-syllabic	Tautosyllabic	
POA of the nasal	Alveolar		Velar
Diphthong	-ia	-ian	-iang
	-ua	-uan	-uang
Monophthong	-a	-an	-ang
	-i	-in	--
	-e	-en	--

In this corpus, the alveolar nasal was either heterosyllabic or tautosyllabic with the preceding vowel, while the velar nasal was always tautosyllabic. The target characters with rimes from Table 1 were put into the carrier sentences in Table 2. Thus,

three types of sentences were constructed: alveolar nasal in heterosyllabic VN sequence /CV₁#nV₂/, alveolar nasal in tautosyllabic VN sequence /CV₁n#V₂/, and velar nasal in tautosyllabic VN sequence /CV₁n#V₂/, Distractor words were also added. Each sentence was repeated twice, and all the prompts were put in a randomized order and then divided into 15 sessions. In this way, the effects of POA, syllabification, vowel type and vowel contexts on VN sequence could all be explored. An example set of sentences with target characters in T1 is shown in Table 2.

Table 2: *Example set of prompt sentences with target segments transcript with IPA and highlighted in boxes (tones ignored)*

/CV ₁ #nV ₂ /					
Characters	她	叫	巴	奶	奶
Pinyin	ta	jiao	ba	nai	nai
IPA	tʰa	teiao	pa	nai	nai
Gloss	She	is called	<name>	grandma.	
/CV ₁ n#V ₂ /					
Characters	她	叫	班	阿	姨
Pinyin	ta	jiao	ban	a	yi
IPA	tʰa	teiao	pan	a	i
Gloss	She	is called	<name>	aunty.	
/CV ₁ n#V ₂ /					
Characters	她	叫	帮	阿	姨
Pinyin	ta	jiao	bang	a	yi
IPA	tʰa	teiao	pan	a	i
Gloss	She	is called	<name>	aunty.	

All the sentences were created in the patterns indicated in Table 2. For the heterosyllabic VN cases (/CV₁#nV₂/), the target character was followed by /nai/; otherwise, it was followed by /a/. The boxed area in Table 2 (the combination of the target character and what follows) built up the *target segment* in this corpus. Thus, the prompt sentences had the same syntactic structure and conveyed very similar semantic meanings. At least 190 tokens per speaker were available for analysis; some speakers had more, because a whole session would be repeated if any of its tokens was mispronounced. The speakers were asked to put emphasis on the <name> of the sentences.

2.3. Real-time MRI, velum opening signal

The real-time MRI data were collected with a 3T MRI system at the Max Planck Institute for Multidisciplinary Sciences in Göttingen, Germany. The data were reconstructed with a frame rate of 50 fps [22]. Noise-suppressing algorithms were implemented in MATLAB to the audio signal which was also time-

The velum opening signal was derived from the MRI images following the procedure in [14]. The time-varying pixel intensities in a region of interest (ROI) were entered as dimensions into principal component analysis (PCA); the scores of the first principal component (PC1) were calculated, resulting in a time varying signal representing the magnitude of velum opening throughout the utterance.

The derived signal is indicative of the degree of velum opening i.e., of nasalization. A larger value indicates a more open velum, and thus larger degree of nasalization. For each speaker, the velum signal was normalized to the range from 0 to 1. Figure 1 (top) shows the variation of velum opening degree across a whole utterance. A rising curve is observed starting from the vocalic segment before the nasal. This curve kept

rising during the VN sequence and reached its maximum later. The yellow triangle in Figure 1 marks the maximum velum opening in the *target segment*. This value of velum opening signal was then extracted for every sentence.

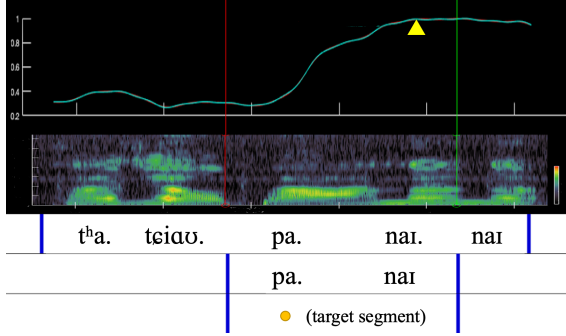


Figure 1: An example of a velum opening signal time-aligned with the spectrogram (top), in which the vertical red/green lines mark the acoustic onset/offset of the target segment; the content of this sentence and the target segment (bottom).

2.4. Segmentation

The audio was segmented into sentences (the first tier in Figure 1). For each sentence, the ‘*target segment*’ was manually decided according to its acoustic boundaries via Praat [23].

2.5. Linear mixed effect modeling

Linear mixed effect models were conducted with the *Pymer* package [24] in Python. The full dataset contained these factors: SPEAKER, GENDER, ONSET (4 levels: /p, p^h, f, Ø/), TONE (4 levels: T1, T2, T3, T4); and properties of the rime: VOWEL_CONTEXT (3 levels: /a, i, e/), VOWEL_TYPE (3 levels: simple, compound-/ia/, and compound-/ua/), POA of the nasal (2 levels: alveolar/velar), SYLLABIFICATION of the VN (2 levels: tautosyllabic/heterosyllabic). Models were built with a full set of fixed factors including VOWEL_CONTEXT, VOWEL_TYPE, POA of the nasal and SYLLABIFICATION and random factors included SPEAKER, ONSET, and TONE.

Model 1 examined the effect of syllabification and vowel context on the maximum velum opening. The dataset contained words with a simple vowel and an alveolar nasal i.e., words with rimes /a/, /an/, /i/, /in/, /e/, /en/ (the boxed area in Table 1). Fixed factors included the VOWEL_CONTEXT and SYLLABIFICATION; while random factors included SPEAKER, ONSET, and TONE. The regressions used random intercepts and slopes.

Model 2 inspected the effect of syllabification, vowel type and POA of the nasal on the maximum velum opening. The full dataset contained target characters in forms of /(V)a/, /(V)an/, and /(V)an/, with V standing for medial /i/ or /u/ i.e., words with rimes /a/, /an/, /aŋ/, /ia/, /ian/, /iaŋ/, /ua/, /uan/, /uaŋ/ (the grayed area in Table 1). The fixed factors included were VOWEL_TYPE, SYLLABIFICATION, and POA of the nasal; and the random factors included were SPEAKER, ONSET, and TONE. The regressions used random intercepts.

3. Results

3.1. Comparison of velum gesture and descriptive analysis

Midsagittal images (top) and the derived velum opening signal (bottom) of a female speaker producing *target segment* /pa#nV₂/, /pan#V₂/, and /paŋ#V₂/ (in T1) are shown in Figure 2 from left to right. The images were extracted at the onset of V₂ to illustrate different degrees of velum opening in each *target segment*, with the velum region highlighted within the yellow bordered square. The lower (more open) the velum, the higher the nasalization degree. In the bottom, the yellow vertical lines mark the onset of the V₂. Commensurate with Figure 2, the velum is lowered (more open) to a higher degree in words with /a, aŋ/ rimes compared to the words with /an/ rime, where the velum is nearly closed, leading to a rime of near-oral quality with nasal coda deleted. Note that in the middle panel (/pan#V₂/), the velum opening signal starts to increase from the beginning of /V₂/). This is due to the velum movement caused by the vowel height change in the /pan#V₂/ sequence, which is induced by nasal loss and vowel raise of the first vowel (/pan/ → /pɛ/).

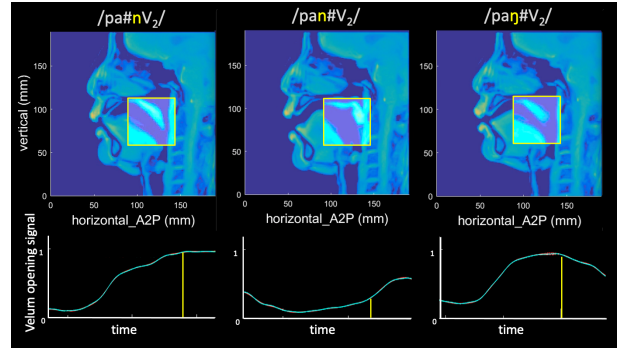


Figure 2: Velum gestures (top) and velum opening signals (bottom) in three target segments; time unit on x-axis: 0.2s; vertical yellow lines mark the onset of V₂.

The maximum of velum opening within the *target segment* was extracted for each repetition, and the values were grouped by rime type. A box plot comparing the maximum velum opening among all rime types is shown in Figure 3. Considerable differences between tokens with /(V)an/ rimes and others are presented: the tokens with rimes /an/, /ian/, and /uan/ have considerably smaller velum opening degrees compared to the other rimes, thereby showing little nasalization.

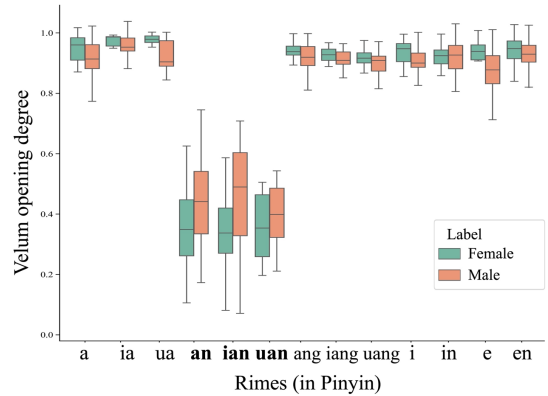


Figure 3: The maximum velum opening in each rime

3.2. Statistic results

Two models were run based on two subsets of the data to explore the effect of different factors influencing the maximum velum opening. Reasons for subsetting the data were the phonological gaps in the Chengdu Dialect which lacks words with velar nasals after /i/ and /e/ (see Table 1). Model 1 was designed based on the boxed rimes in Table 1 to explore the influences of SYLLABIFICATION and VOWEL_CONTEXT, and therefore did not contain the rimes with diphthong and velar nasal. Model 2 was designed based on the grayed rimes in Table 1 to investigate the effects of SYLLABIFICATION, POA, and VOWEL_TYPE and therefore did not contain the rimes that lack velar nasal counterparts. To investigate each effect in model 2, rimes that lack the expected counterparts were excluded, e.g., when investigating the effect of SYLLABIFICATION, the rimes with tautosyllabic velar nasals were dropped because heterosyllabic velar nasal counterparts did not exist.

In model 1, LME regressions with pair-wise comparisons of words with tautosyllabic VN sequence showed significantly smaller velum opening in words with /an/ rime than in words with other rimes (Table 3). In the comparison of words with heterosyllabic V#N sequence, no such significant difference was found, regardless of the vowel context. When comparing the effect of syllabification within the same vowel context, the only significant effect emerged between word pairs with /a/ and /an/ rimes, but not in word pairs with /i/-/in/ or /e/-/en/ rimes.

Table 3: Results of model 1 for the effects of vowel context (-a, -i, -e) and syllabification (tauto-/heterosyllabic) on maximum of velum opening.

Fixed factor	Compare between	p-value	
VOWEL_CONTEXT	-an vs. -in	$p < 0.001$	***
	-an vs. -en	$p < 0.001$	***
	-in vs. -en	$p = 0.641$	
SYLLABIFICATION	-a vs. -an	$p < 0.1$	*
	-e vs. -en	$p = 0.998$	
	-i vs. -in	$p = 0.952$	

In model 2 we compared the maximum of velum opening between words with different syllabifications, POAs, and vowel types. Maximum velum opening was affected significantly by syllabification (/an, ian, uan/ - /a, ia, ua/) and POA (/an, ian, uan/ - /aŋ, iaŋ, uaŋ/). Words with /(V)an/ rimes were characterized by a significantly smaller velum opening than words with /(V)a, (V)aŋ/ rimes (Table 4, row SYLLABIFICATION and POA). Unlike syllabification and POA, VOWEL_TYPE did not have a significant influence on the full dataset or on words with /(V)an/ rimes. However, in words with /(V)a/ rimes, there was a significantly greater degree of maximum velum opening in /ua/ than in /a, ia/ rimes, and in words with /(V)aŋ/ rimes, the maximum velum opening was significantly smaller in /uaŋ/ than in /aŋ/ rimes (Table 4).

In conclusion, it was tested that the maximum velum opening was significantly smaller in words with a tautosyllabic VN sequence where V is a low vowel and N is alveolar i.e., in words with /(V)an/ rimes. Significant differences among vowel types were only observed in /(V)a, (V)aŋ/ rimes, but not in /(V)an/ rimes, where the nasal was lost. Descriptively, the maximum velum opening for female speakers was larger in all rimes but /(V)an/, where the sound change took place.

Table 4: Results of model 2 for the effects of syllabification (tauto-/heterosyllabic), POA of the nasal (alveolar/velar), and vowel type (simple/compound vowel) on maximum velum opening.

Fixed factor	Compare between			p-value	
SYLLABIFICATION	/Vn/	vs.	/V#n/	$p < 0.001$	***
POA	/Vn/	vs.	/Vŋ/	$p < 0.001$	***
	/a/	vs.	/ia/	$p = 0.279$	
VOWEL_TYPE in /(V)a#n/ rimes	/a/	vs.	/ua/	$p < 0.001$	***
	/ia/	vs.	/ua/	$p < 0.001$	***
VOWEL_TYPE in /(V)aŋ/ rimes	/aŋ/	vs.	/iaŋ/	$p = 0.143$	
	/aŋ/	vs.	/uaŋ/	$p < 0.1$	*
	/iaŋ/	vs.	/uaŋ/	$p = 0.260$	

4. Discussion

The results of this study show that, in varies VN sequences in the Chengdu Dialect, the maximum velum opening is significantly smaller in tokens with /(V)an/ rimes than tokens with the heterosyllabic VN sequence /(V)a#n/, a velar nasal /(V)aŋ/, and the other vowel contexts /in, en/.

The maximum velum opening being affected by syllabification (/an/ vs. /a#n/), vowel context (/an/ vs. /in, en/) and the place of articulation of the nasal (/an/ vs. /aŋ/) suggests a phonetically motivated nasal loss like that in American English [12], Camuno Italian [13] and northern Italian dialects [25]. Different from the above languages/dialects, this sound change in the Chengdu Dialect only takes place in words with /(V)an/ rimes. In addition, this nasal loss pattern was unaffected by vowel type in /(V)an/ rimes, but the degree of nasalization is significantly larger in /ua#n/ than in /a#n, ia#n/, while significantly smaller in /uaŋ/ than in /aŋ/.

This study provides empirical evidence for the diachronic nasal loss in /(V)an/ rime words of the Chengdu Dialect. During this sound change, the vowels have lost coarticulatory traces of nasality and have become plain oral vowels. The nasal loss together with the observed diachronic vowel raising in /(V)an/ syllables [17] are consistent with other research showing low vowel raising before nasal consonants in other languages [15], [16]. As such, the present data not only provides evidence for pathway B (see 1. Introduction) through which vowel nasalization disappears, but also specifies the co-occurrence of vowel raising and nasal loss in this dialect. In the future, more speakers will be recruited to investigate the interaction among different articulators involved in these VN sequences to gain a more integrated view on the vowel raising and nasal loss in the Chengdu Dialect.

5. Acknowledgements

The first author is funded by the China Scholarship Council (CSC) (grant No. 202008440473). This research is funded by the project InterAccent, which has received funding from the European Research Council (ERC) under the European Union's Horizon 2020 research and innovation program (grant agreement No. 742289). Website: <https://www.phonetik.uni-muenchen.de/Forschung/interaccent/interAccent.html>

6. References

- [1] S. Zhen, “Chengdu Yuyin de Chubu Yanjiu (A preliminary study on Chengdu phonetics),” *J. Sichuan Univ. Soc. Sci. Ed.*, no. 1, 1958.
- [2] S. Zhen, “Sichuan Fangyan de Biweiyun (Nasal codas in Sichuan dialect),” *Fangyan (Dialect)*, no. 4, pp. 241–243, Nov. 1983.
- [3] M. Y. Chen, “Acoustic analysis of simple vowels preceding a nasal in Standard Chinese,” *J. Phon.*, vol. 28, no. 1, pp. 43–67, 2000, doi: 10.1006/JPHO.2000.0106.
- [4] X. Shi, J. Zhang, and F. Shi, “Influential factors on nasal codas in Mandarin Chinese: evidence from phonetic experiments,” *Stud. Chinese Lang.*, vol. 5, pp. 578–589, 2019.
- [5] J. Cui, Y. Chen, and L. Wang, “Wuxi Speakers’ Production and Perception of Coda Nasals in Mandarin,” in *Interspeech*, 2018, pp. 2559–2562, doi: 10.21437/Interspeech.2018-2224.
- [6] M. Faytak, S. Liu, and M. Sundara, “Nasal coda neutralization in Shanghai Mandarin: Articulatory and perceptual evidence,” *Lab. Phonol.*, vol. 11, no. 1, pp. 1–29, Dec. 2020, doi: 10.5334/LABPHON.269.
- [7] T. Baer, J. C. Gore, L. C. Graeco, and P. W. Nye, “Analysis of vocal tract shape and dimensions using magnetic resonance imaging: vowels,” *J. Acoust. Soc. Am.*, vol. 90, no. 2 Pt 1, pp. 799–828, 1991, doi: 10.1121/1.401949.
- [8] K. Mády *et al.*, “Use of real-time MRI in assessment of consonant articulation before and after tongue surgery and tongue reconstruction,” in *Proc. 4th International Speech Motor Conference*, 2001, pp. 142–145.
- [9] C. Carignan *et al.*, “Analyzing speech in both time and space: Generalized additive mixed models can uncover systematic patterns of variation in vocal tract shape in real-time MRI,” *Lab. Phonol.*, vol. 11, no. 1, pp. 1–26, Mar. 2020, doi: 10.5334/LABPHON.214.
- [10] C. Carignan *et al.*, “Planting the seed for sound change: Evidence from real-time mri of velum kinematics in German,” *Language (Baltim.)*, vol. 97, no. 2, pp. 333–364, 2021, doi: 10.1353/LAN.2021.0020.
- [11] P. S. Beddor, “A coarticulatory path to sound change,” *Language (Baltim.)*, vol. 85, no. 4, 2009, doi: 10.1353/lan.0.0165.
- [12] J. J. Ohala and M. G. Busa, “Nasal loss before voiceless fricatives: a perceptually-based sound change,” *Riv. Linguist.*, vol. 7, pp. 125–144, 1995.
- [13] M. Cresci, “VN > V in Camuno: An alternative historical pathway to nasal loss,” *Ital. J. Linguist.*, vol. 31, no. 1, pp. 61–92, 2019, doi: 10.26346/1120-2726-132.
- [14] C. Carignan *et al.*, “The phonetic basis of phonological vowel nasality: Evidence from real-time MRI velum movement in German,” in *Proceedings of 19th International Congress of Phonetic Sciences (ICPhS)*, 2019, pp. 413–417.
- [15] M. Chen, “An Areal Study of Nasalization in Chinese,” *J. Chinese Linguist.*, vol. 3, pp. 16–59, 1975.
- [16] J. Hajek, *Universals of Sound Change in Nasalization*. Oxford: Blackwell Publishers, 1997.
- [17] H. Hu and Y. Zhang, “Path of Vowel Raising in the Chengdu Dialect of Mandarin,” in *Proceedings of the 29th North American Conference on Chinese Linguistics (NACCL-29)*, 2017, vol. 2, pp. 481–498.
- [18] A. Grainger, *Xishu Fangyan (Western Mandarin)*. Shanghai: American Presbyterian Mission Press, 1900.
- [19] S.-L. Zhen, “Sichuan fangyan yinxi (The phonetic system of the Sichuan dialect),” *J. Sichuan Univ. (Social Sci. Ed.)*, 1960.
- [20] Y. Zhang, “An Analysis of Synchronic Variations in the Phonetic System of Chengdu-dialect,” *Hanyu Xuebao*, no. 3, pp. 39–45, 2010.
- [21] Dialect Research Group of Sichuan University, “The Phonology of Sichuan Dialect,” *J. Sichuan Univ. Soc. Sci. Ed.*, no. 3, 1960.
- [22] M. Uecker, S. Zhang, D. Voit, A. Karaus, K. D. Merboldt, and J. Frahm, “Real-time MRI at a resolution of 20 ms,” *NMR Biomed.*, vol. 23, no. 8, pp. 986–994, Oct. 2010, doi: 10.1002/NBM.1585.
- [23] P. Boersma and D. Weenink, “Praat: doing phonetics by computer.” 2022, [Online]. Available: <http://www.praat.org/>.
- [24] E. Jolly, “Pymer4: Connecting R and Python for Linear Mixed Modeling,” *J. Open Source Softw.*, vol. 3, no. 31, p. 862, Nov. 2018, doi: 10.21105/JOSS.00862.
- [25] J. Hajek and S. Maeda, “Investigating universals of sound change: the effect of vowel height and duration on the development of distinctive nasalization,” in *Papers in laboratory phonology V*, M. Broe and J. Pierrehumbert, Eds. Cambridge: Cambridge University Press., 2000, pp. 52–69.