

한국인 스페인어 학습자의 음향음성학적 모음 발음 연구*

김 주 경
서울대학교

김주경(2020), 「한국인 스페인어 학습자의 음향음성학적 모음 발음 연구」, 이베로아메리카 연구, 31(1), 23-52.

초 록 본 논문은 실험음성학적 방법을 도입하여 한국어와 스페인어 단모음 체계를 비교하고, 한국인 스페인어 학습자의 모음 발음을 과학적으로 분석하는 것을 목적으로 하였다. Flege의 음성 학습모델(Speech Learning Model) (Flege 1995)을 적용하여, 스페인어 모음을 한국어 모음 유사도에 따라 4단계로 분류한 후 스페인어 모음이 한국어와 유사할수록 습득하는 데 도움이 되는지 알아보았다. 그리고 학습자 그룹을 학습 경험에 따라 두 개로 나누어 원어민 그룹과 비교함으로써 스페인어 경험량 증가에 따라 모음 발음이 원어민 수준에 다다르게 되는지 파악하였다. 우선 스페인어 모어 화자의 청취실험을 통해 실제로 한국인 학습자의 비원어민적인 발음으로 인식 오류가 발생할 수 있음을 확인하였다. 그리고 산출실험을 통해 원어민과 학습자가 실제로 어느 위치에서 스페인어 모음을 조음하는지를 알아보았다. 총 24명의 실험 참여자에게 45개의 스페인어 문장을 낭독하도록 하였고, 녹음한 자료는 프랏(Praat)으로 모음 구간 중간 지점의 F1, F2를 측정하였다. 측정값을 분석한 결과 한국어와 유사도가 높은 모음들은(/a/, /e/, /i/, /u/) 원어민과 학습자 간 차이가 크게 드러나지 않은 반면, /o/는 유의미한 차이가 있었다. 하지만 학습량에 따른 차이는 없는 것으로 나타났다. 이러한 현상은 실제로는 두 모음 체계 간 유의미한 차이가 있음에도 직관적으로 한국어 모음을 스페인어 모음에 일대일 대응시켜 인식하고 발음하도록 하는 교수법과 학습법에 기인한 것으로 생각된다.

핵심어 L2 스페인어, 모음 습득, 모음 산출, 발음 오류, 실험음성학

* 본 논문은 저자의 석사학위논문(서울대학교, 2020)을 기반으로 작성되었습니다.

I. 서론

최근 외국어 소통 능력이 강조됨에 따라 외국어 음소를 정확하게 인식하고 발음하는 능력에 대한 많은 연구가 진행되었다. 특히 알파벳을 사용하는 언어(영어, 스페인어, 포르투갈어, 프랑스어, 독일어 등) 간 연구는 이미 상당하며, 영어와 한국어를 비교하는 연구도 다수 진행되었다(Flege *et al.* 1997; 윤영도 2007; Lee and Iverson 2012). 하지만 아직 한국어와 스페인어 간 연구는 미비한 편이며, 모음 습득 연구는 사실상 거의 없다고 해도 과언이 아니다(자모음 연구: 강현화·신자영·이재성·임효상 2003; 서경석 2007; 김승한 2009; 김원필 2009, 2013, 2014; 자음 연구: 김원필 2001, 2017; 문정 2014; 모음 연구: 이신영 2015).¹⁾

이러한 한국인 스페인어 학습자를 대상으로 한 음소 습득 연구의 부족은 스페인어 발음이 용이하다는 사고가 학습자뿐만 아니라 교수자 사이에서도 만연하기 때문으로 보인다. 이러한 인식 때문인지 최근까지도 한국인 스페인어 학습자의 분절음 습득, 특히 그중에서도 모음은 스페인어학 연구 분야에서 주요 관심사가 아니었다. 비록 김원필(2009, 2013), 서경석(2007) 등이 스페인어와 한국어 모음 체계를 대조하고 그 차이를 제시하였지만 과학적인 분석을 통해서가 아닌 연구자의 직관과 청각에만 의존한 조음음성학적 탐구였다는 점에서 한계를 가진다. 사실 한국어가 스페인어에 비하면 훨씬 더 많은 수의 모음과 넓은 모음 공간(vowel space)을 가지기 때문에 스페인어 학습자가 한국어를 배울 때 그 반대의 경우보다 모음 습득이 더 어려울 것이라는 생각이 자연스럽다. 한국어 단모음은 최소 7개(/a, e, i, o, u, ʌ, ɯ/, (Shin *et al.* 2012))인 반면 스페인어는 5개고(/a, e, i, o, u/ (Hualde 2005)), 이 5개 모두 한국어 모음으로 대응될 수 있는 것처럼 보인다. 실제로 많은 스페인어 학습 교재에서 용이한 학습을 위해 스페인어 발음을 한국어 모음에 일대일 대응시켜 한글로 표기한다. 그러나

1) 이중 서경석(2007), 김원필(2009), 이신영(2015)는 스페인어 모어 화자를 대상으로 진행한 연구다.

기존 인식과는 달리 스페인어 모음의 조음 위치는 엄밀히 따지면 한국어와 다르며, 본고에서는 이러한 차이에서 비롯된 한국인 스페인어 학습자의 모음 발음 오류를 실험음성학적으로 제시하고자 한다.

II. 연구 배경

1. Flege의 음성학습모델(Speech Learning Model)

본 연구는 L2 학습자가 L1 음소 범주를 토대로 인식 및 발음한다는 주장을 토대로 한 Flege의 음성학습모델(Speech Learning Model; 이하 SLM) (Flege 1987, 1992, 1995; Bohn and Flege 1990, 1992)을 도입하여 L1 한국어-L2 스페인어 모음을 분석하려고 한다. SLM에 따르면 학습자가 본인의 모국어 음소와 유사한 L2 음소 간 차이를 인식할 수 없을 때 그 학습자의 해당 음소 발음은 정확할 수 없다. 즉, L2 음소 인식 오류는 발음 오류로 이어질 수 있다는 것이다. 이러한 발음 오류는 결정적 시기가 지난 성인 학습자들이 고치기 힘든 것으로 흔히 알려져 있다. 하지만 SLM은 성인 학습자라도 학습경험 정도에 따라 L2 음소 인식과 발음이 개선될 수 있다고 보았다(학습효과, Learning Effect) (Flege 1987; Flege *et al.* 1997). 본고에서는 L2 학습자들이 경험치가 쌓이면서 L1과 L2 모음 간 음성적 차이를 구분할 수 있게 되며, 장기적으로 원어민과 유사하게 인식 및 발음하게 된다는 가설을 한국인 스페인어 학습자를 대상으로 확인해 보고자 한다.

그리고 SLM에서는 L1과 L2 간 음소 대응 패턴에 따라 음소를 세 가지 그룹, 즉 “동일한 소리(identical phones)”, “유사한 소리(similar phones)”, “새로운 소리(new phones)”로 분류한다. L1과 L2 음소가 음성적으로 동일한 특성을 가질 경우 “동일한 소리”²⁾, 같은 IPA 기호를 갖지만(혹은 그럴듯한 L1-L2 대응이 있지만) 어느 정도 차이가 있을 경우 “유사한 소리”³⁾로 분류하였다. 그리

2) 한국어와 영어의 /m/를 예로 들 수 있다.

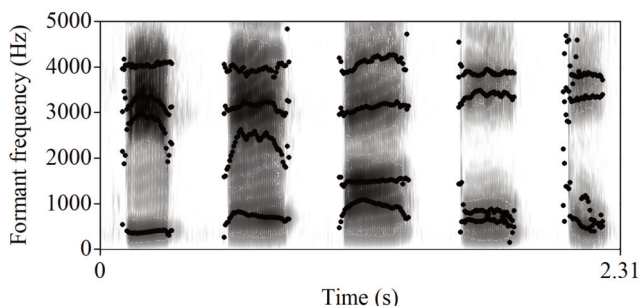
3) 프랑스어와 영어의 /t/를 예로 들 수 있는데, 프랑스어에서 /t/는 무성 치경 파열음으로, 영어는 유성 경구개 파열음으로 실현된다(Flege 1987, 48).

고 L1 대응이 없는 L2 음소를 “새로운 소리”⁴⁾라고 보았다. 이때 습득에 문제가 되는 소리는 유사한 소리와 새로운 소리인데, SLM에 따르면 학습자 입장에서는 동등분류 현상(Equivalence Classification)으로 인해 유사한 소리가 새로운 소리보다도 습득하기 더 어려울 수 있다고 한다. 학습자가 L1과 유사한 L2 소리를 있는 그대로 인식하지 않고, 동일하지 않더라도 비슷한 L1 소리에 대응시켜 버려(Flege 1987) 유사한 소리의 경우 그 차이를 인지하는 데 어려움을 겪는다는 것이다.

본고에서는 이러한 이론적 배경을 바탕으로, 스페인어 모음을 한국어 유사 정도에 따라 분류한 후 유사 정도와 학습 경험량에 따른 학습자의 모음 실현 양상을 실험을 통해 알아보고자 한다.

2. 한국어-스페인어 단모음 체계 대조

두 언어체계 간 모음 발음을 비교하기 위해서는 각 모음이 조음될 때 혀의 위치와 입술의 원순성 정도를 알아야 한다. 입의 모양 변화에 따라 강조되는 소리의 주파수가 변하여 각 모음이 고유한 음가를 가지게 되기 때문이다. 이러한 모음별 특징은 스펙트로그램 상 포먼트로 나타나기 때문에 음향분석 프로그램인 프랏(Praat)을 이용, 스펙트로그램을 그려 모음을 분석할 수 있다. <그림 1>



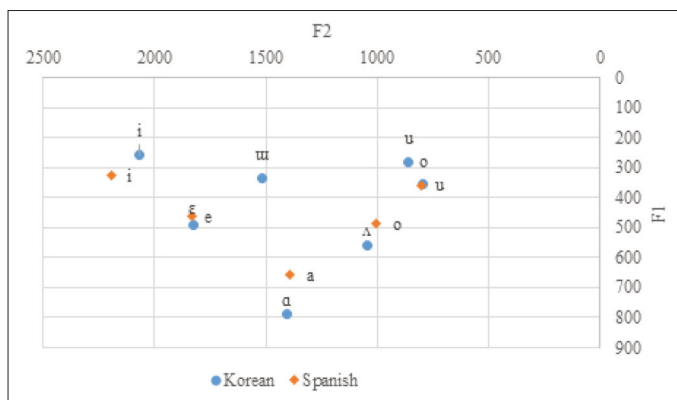
<그림 1> /pipepapopu/를 발음한 스펙트로그램

4) 한국어를 학습하는 스페인 원어민의 입장에서 /ㅡ/와 같은 소리가 이 경우에 해당한다.

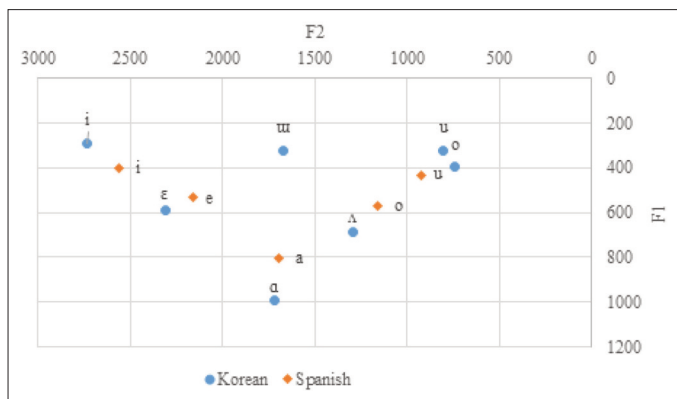
에서 진한 색의 띠가 두텁게 나타나는 부분이 포먼트인데, 모음마다 그 형태가 전혀 다르게 나타나는 것을 확인할 수 있다(앞에서부터 /i, e, a, o, u/). 검은 점으로 표시된 가장 아래에 위치한 띠부터 순서대로 F1, F2, F3에 해당한다. 예를 들어, /i/는 F1이 낮고, F2가 높이 위치함을 확인할 수 있다. 대체로 혀의 위치에 따라 F1과 F2가, 입술 모양에 따라 F3가 결정된다.

그런데 대부분의 모음, 특히 스페인어 단모음은 첫 두 포먼트(F1, F2)로 구분할 수 있기(Catford 2001; Clegg and Fails 2018 등) 때문에 본 연구에서도 F1과 F2만 분석 대상에 포함한다. F1값은 모음의 높이(height)와 관련 있는 수치로, 혀의 높이에 반비례하며, F2는 전설성(frontness)과 관련하여 혀가 앞에 놓일수록 값이 높아진다. 즉, 고모음(high vowel)일수록 F1값이 낮아지고 저모음(low vowel)일 때 높아진다. 그리고 전설모음(front vowel)에 가까울수록 F2값이 높고 후설모음(back vowel)인 경우 그 값이 낮아진다. 예를 들어, 전설 고모음인 [i]의 경우 중설 저모음인 [a]에 비해 F1값이 낮고 F2값이 높을 것으로 예상할 수 있다.

본격적인 실험에 앞서 포먼트 수치를 기반으로 한국어와 스페인어 모음체계를 비교분석 해본다. 한국어와 스페인어 모음 유사 정도를 판단하기 위한 포먼트 값 자료로 각각 Shin *et al.*(2012)과 Chládková *et al.*(2011)의 연구를 활용하였다. 스페인어 모음이 한국어 모음에 일대일 대응된다는 대부분 한국인 학습자들의 인식과 달리 실제 F1, F2 수치를 이용하여 그린 F1-F2 평면도를 살펴보면 서로 일대일 대응시키는 것이 쉽지 않다(<그림 2>, <그림 3>). 전반적으로 스페인어(마름모)보다 한국어(원)에서, 특히 여성 화자들이 포먼트 평면도를 비교적 넓게 사용한다. 개별적인 모음을 살펴보면, /i/보다 /ɪ/가 더 위에서, /a/보다 /ɒ/가 더 아래에서 조음되며, /o/에 /ɔ/보다 /ɯ/가 더 가까울 뿐만 아니라 /u/에 가까운 모음은 /ʊ/보다 /ɤ/다. 이러한 경향은 남성 화자들에서 더 강하게 나타난다. 스페인어 /u/는 한국어 /ɤ/와 조음위치가 거의 일치한다.



〈그림 2〉 한국어-스페인어 모음평면도(남성화자)



〈그림 3〉 한국어-스페인어 모음평면도(여성화자)

이를 수치화해서 제시하면 두 체계 간 차이가 조금 더 명확하게 드러난다. 평면도에서 유클리드 거리 공식(ED; Euclidean Distance)을 사용하여 모음 간 거리를 계산해 보면 /o/-/Λ/, /u/-/ɿ/ 가 일대일 대응쌍으로 나타나지 않는다. /o/는 성별과 관계없이 /ɿ/가 가장 가까운 한국어 모음으로 나타난다 (/ɿ/-/o/ (ED_M=83.90; ED_F=183.51); /Λ/-/o/ (ED_M=246.24; ED_F=448.71)). /u/는 성별에 따라 약간의 차이가 있는데, 남성의 경우 /Λ/와 거의 동일한 조음위치를 가지는 반면, 여성화자는 미세한 차이지만 /ɿ/와 더 유사한 것으로

보인다(/ ㅏ /- /u/ (ED_M=6.25, ED_F=184.26); / ㅓ /- /u/ (ED_M=100.37, ED_F=162.94)).

앞서 모음평면도를 그려본 결과, 스페인어 모음을 한국어 모음에 일대일 대응시킬 수 없으며, SLM을 적용해 보았을 때 단순하게 같은 소리-유사한 소리-새로운 소리의 세 가지로 분류하기가 곤란함을 알 수 있다. 여러 연구(Flege 1987; Flege and Hillenbrand 1984; Yang 1996; 김주연 2014 등)에서는 이 분류체계를 사용하였지만 지금까지도 ‘유사함’의 유무에 대한 구체적인 기준이 없다는 점을 고려할 때 3분류 체제는 보완이 필요하다. 따라서 본고에서는 권성미(2009)처럼 유사성이 ‘정도’에 따라 연속선상에 있는 것으로 보았고, 다음과 같이 네 단계로 분류하였다. 사실상 같은 소리라고 볼 수 있는 / ㅐ /- /e/ 쌍은 가장 높은 유사4단계 소리로, / ㅣ /- /i/ 쌍, / ㅓ /- /a/ 쌍은 평면도 상 어느 정도 거리감이 있지만 /u/나 /o/에 비해서는 명확하게 일대일 대응쌍으로 나타나기 때문에 유사3단계로 분류하였다. 한편, /u/는 / ㅓ /랑 가깝긴 하지만 / ㅏ /와도 매우 가깝게 위치하기 때문에 한국인 학습자 입장에서 어느 정도 혼란을 초래할 수 있을 것으로 예상하여 유사2단계로 보았다. 마지막으로, /o/는 앞서 설명했듯이 / ㅏ /보다도 / ㅣ /와 유사하게 발음될 정도로 유사하지 않은 대응쌍이기 때문에 유사1단계 소리로 분류하였다.

- (i) 유사4단계 소리(같은 소리) : / ㅐ /- /e/
/ ㅐ /- /e/ (ED_M=25.80; ED_F=161.23)
- (ii) 유사3단계 소리 : / ㅣ /- /i/, / ㅓ /- /a/
/ ㅣ /- /i/ (ED_M=145.96; ED_F=202.43)
/ ㅓ /- /a/ (ED_M=131.33; ED_F=191.56)
- (iii) 유사2단계 소리 : / ㅏ /- /ㅓ /- /u/
/ ㅓ /- /u/ (ED_M=100.37, ED_F=162.94)
- (iv) 유사1단계 소리 : / ㅣ /- /ㅏ /- /ㅓ /- /o/
/ ㅏ /- /o/ (ED_M=246.24; ED_F=448.71)

3. 한국인 학습자의 모음 발음 오류 현상

앞서 우리는 한국어 단모음이 스페인어와 어느 정도 차이가 있음을 확인하였다. 이 장에서는 한국인 스페인어 학습자의 모음 발음을 원어민이 들었을 때 얼마나 정확하게 인식하는지 간단한 실험을 통해 알아보고자 한다.

1) 인식실험 자료 및 진행 방법

한국인 스페인어 학습자 4명(남 2명, 여 2명)의 음성 녹음 파일을 잘라 본 실험의 자료로 활용하였다. 특정 단어가 연상될 수 없도록 C(자음)+V(모음)+C(자음)에 해당하는 부분을 잘라서 사용하였고,⁵⁾ 발음 상 문제가 예상되는 부분이 /o/와 /u/기 때문에 /a/, /e/, /i/는 각 4번만, /o/, /u/는 각 10번씩 본 실험 자료로 활용되었다. 총 32개의 소리 파일을 임의의 순서로 4명의 스페인어 모어 화자에게 들려주고 V만 다르게 5지선다로 문항지를 구성하여 1개의 응답을 선택하게 하였다. 예를 들자면 ‘pop’을 들려주고 ‘pap – pep – pip – pop – pup’ 중 하나를 선택하게 했다는 것이다. 본격적인 실험 전 예시로 소리 하나를 들려주어 진행 방법에 대한 의문점이 없도록 하였고, 피험자가 못 들었다고 할 경우 다시 한 번 들려주었다.

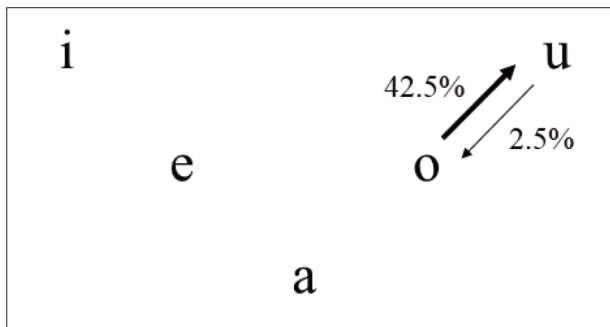
2) 인식실험 결과

실험 실시 후 각 문항에 대한 정답률을 계산하였다. 예상했던 대로 /a/, /e/, /i/ 발음에 대한 정답률은 100%에 달했다. 각 모음에 대해 네 문항씩 출제하였는데 모든 원어민 피험자가 모든 문항에 대해 올바르게 답하였다.

하지만 /o/와 /u/의 경우 어느 정도의 인식 오류가 발생하였다(<그림 4>). 전체 32문항 중 10문항이 /u/로 출제되었는데, 한 명의 피험자가 한 개의 문항에서 잘못 응답하였다(총 정답률 97.5%(=39/40)). 실제 의도한 발음은 ‘pup’이었는데 ‘pop’으로 응답한 경우였다. 한편, /o/도 /u/처럼 10문항이 출제되었는데, 오류 비율이 훨씬 높게 나타났다(총 정답률 42.5%(=17/40)). 한국인

5) 자음의 경우 스펙트로그램에서 관찰 대상인 모음과의 경계가 비교적 분명하게 나타나는 /p/, /t/ 혹은 /k/로, 모음은 강세 오는 모음으로 구성하였다.

학습자가 발음한 /o/에 대해 절반 이상의 경우 /u/로 인식됐다는 것이다. 이러한 오류는 발화자나 청취자 모두에게서 공통적으로 발견되어, 개별 실험 참여자의 문제라고 볼 수 없었다. 또한, 앞뒤 자음(C)과 관계없이, 고르게 오류가 발생했기 때문에 특정 자음의 영향도 아닌 것으로 나타났다. 이 실험 결과는 앞서 <그림 2>, <그림 3>의 모음평면도를 통해 어느 정도 해명이 가능하다. 한국어 /ㅜ/, /ㅡ/ 둘 다 스페인어 /u/ 근처에서 조음되는 것으로 보이는데, 한국인 학습자의 L1이 스페인어로 전이된다고 하면 스페인어 모어 화자가 /o/를 자주 /u/로 인식하는 현상을 설명할 수 있다. 그리고 이 결과는 한국인의 /o/와 /u/ 발음이 구별이 어려워 스페인어 의사소통 시 오해의 소지가 있을 수 있음을 시사한다.⁶⁾ 스페인어 모어 화자와 다음 장에서 다룰 모음 발음 실험을 통해 한국인 학습자와 스페인 모어 화자의 스페인어 모음 발음이 실제로 어떤 차이를 갖는지 실험음성학적으로 증명해 보려고 한다.



<그림 4> 인식 실험 인식 오류 분포⁷⁾

- 6) 이처럼 /o/ 인식 오류가 높음에도 지금까지 이에 대해 진행된 연구가 없었던 것은 아마 평소 대화 상황에서는 문맥이 있기 때문에 실질적으로 소통하는 데 큰 문제가 없었기 때문인 것으로 생각된다. 하지만 본 인식실험 결과에 따르면 문맥이 없을 경우 의사소통 시 오해가 발생할 가능성이 있다.
- 7) Morrison(2003) 연구에서 사용한 그래프를 참고하였다. /i/, /e/, /a/의 경우 인식 오류가 발생하지 않아 아무런 표시가 없고, 화살표가 굵을수록 인식 오류가 빈번하게 일어났음을 의미한다.

III. 모음 산출실험 절차

1. 산출실험 참여자

스페인어 모음 산출실험에는 스페인어권 원어민 4인과 한국인 학습자 20인이 참여하였다. 모든 참여자는 만 18세 이상의 대학 재학생 혹은 졸업생이며, 표준 스페인어 혹은 한국어를 모국어로 구사하는 사람들이었다. 원어민(이하 NS)⁸⁾ 그룹과의 모음 인식 및 발음 차이가 학습 경험에 따라 달라지는지 살펴 보기 위해 한국인 학습자는 두 그룹으로 나누었다. 고급 학습자(이하 KSA) 그룹에는 5년 이상 학습 경험이 있는 참여자를, 초급 학습자(이하 KSB) 그룹에는 3개월 이상이면서 1년 미만의 학습 경험이 있는 참여자를 포함하여⁹⁾ 그룹 간 수준 차이를 명확하게 하였다. KSA는 12명, KSB는 8명이었다. 구체적인 피험자 정보는 다음 <표 1>, <표 2>, <표 3>과 같다. 세 그룹 내 성별 분포는 제각각이었기 때문에 성별에 따른 포먼트 값 차이는 정규화 작업을 통해 조정하였다.

<표 1> 스페인어 모어 화자 정보(NS)

실험자	성별	출신지	한국어 학습 시작 나이	한국어 학습 기간	한국 체류 경험
NS1	여	스페인	22	1~3년 사이	5~10년 사이
NS2	여	멕시코	30	1년 이내	3~5년 사이
NS3	여	멕시코	33	1~3년 사이	5~10년 사이
NS4	여	스페인	해당 없음	해당 없음	해당 없음

8) NS: Native Speaker, KSA: Korean Speaker - Advanced, KSB: Korean Speaker - Beginner

9) KSA들은 대부분 스페인어를 전공하는 졸업반 대학생들 혹은 대학원생들로, 대부분 스페인어권 거주 경험이 있었다. KSB 그룹 피험자는 모두 서울대학교에서 개설하는 초급스페인어1을 수강 완료한 학생들이었으며, 대부분 실험 참가 시점에 초급스페인어2를 수강 중이었다. 최대한 두 학습자 그룹 간 경험량 및 수준 차이가 극명하게 나도록 학습자 그룹을 설정하였다.

〈표 2〉 스페인어 고급 학습자 정보(KSA)

실험자	성별	출신지	모국어	스페인어 학습 시작 나이(만)	스페인어 학습 기간	스페인어권 체류 경험	타 외국어 학습 경험
KSA1	여	대전	한국어	16	10년 이상	스페인 (10개월)	일본어 (1년)
KSA2	남	서울	한국어	16	5-10년 사이	스페인 (6개월)	-
KSA3	여	경기	한국어	19	10년 이상	스페인 (6개월)	-
KSA4	여	서울	한국어	15	5-10년 사이	스페인 (6개월)	포르투갈어 (1년)
KSA5	여	서울	한국어	15	5-10년 사이	스페인 (12개월)	중국어 (4년)
KSA6	남	서울	한국어	16	5-10년 사이	-	-
KSA7	여	대전	한국어	16	5-10년 사이	-	포르투갈어 (6개월)
KSA8	여	대전	한국어	18	5-10년 사이	스페인 (5개월)	포르투갈어 (6개월)
KSA9	남	서울	한국어	18	5-10년 사이	멕시코 (5개월)	-
KSA10	남	서울	한국어	16	5-10년 사이	스페인 (6개월)	프랑스어 (1년)
KSA11	여	서울	한국어	16	5-10년 사이	스페인 (6개월)	중국어(1.5년), 일본어(3년)
KSA12	남	서울	한국어	19	4-5년 사이	스페인 (6개월)	-

〈표 3〉 스페인어 초급 학습자 정보(KSB)¹⁰⁾

실험자	성별	출신지	모국어	스페인어 학습 시작 나이(만)	스페인어 학습 기간	타 외국어 학습 경험
KSB1	여	서울	한국어	18	1년 미만	프랑스어(3년)
KSB2	남	경기	한국어	19	1년 미만	중국어(1년)
KSB3	남	대전	한국어	19	1년 미만	-
KSB4	남	전주	한국어	20	1년 미만	-
KSB8	남	광주	한국어	19	1년 미만	-
KSB9	남	서울	한국어	19	1년 미만	-
KSB10	여	강릉	한국어	18	1년 미만	-
KSB11	남	전주	한국어	21	1년 미만	-

2. 산출실험 자료

산출실험 자료로 스페인어로 된 45개의 단순한 문장을 사용하였다. 실험 문장은 Chládková *et al.*(2011)의 연구에서 사용한 문장들 “①En ②C₁VC₂e ③y ④C₁VC₂o ⑤tenemos ⑥V.”에서 두 번째, 네 번째, 여섯 번째 단어를 바꾸어가면서 구성하였다. 문장 틀에서 C (Consonante)는 /p, t, k/를 조합하였고, V (Vocal)는 /a, e, i, o, u/를 조합하였다. 그리고 같은 문자에는 동일한 소리를 부여하였다. 즉, C₁는 C₁끼리, V는 V끼리, C₂는 C₂끼리 같은 글자를 가지게 되어 아래와 같은 실험 문장들이 생성된다. 이렇게 자음 3개와 모음 5개의 모든 조합으로 문장을 구성하면 총 45개(자음 3개×모음 5개×자음 3개)가 만들어진다. 이렇게 구성된 실험 문장에서 ②와 ④의 V로 표시된 모음이 본 실험의 분석 대상이 되었다.

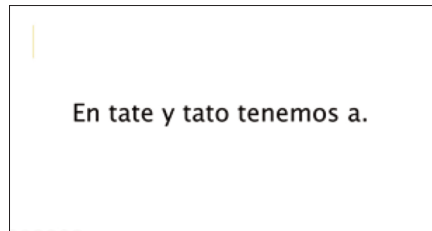
〈표 4〉 산출실험 문장 예시

En <u>pate</u> y <u>pato</u> tenemos a.
En <u>quete</u> y <u>queto</u> tenemos e. ¹⁾
En <u>tope</u> y <u>topo</u> tenemos o.

3. 산출실험 방법

PowerPoint 프레젠테이션을 이용하여 참여자에게 한 슬라이드에 하나의 실험 문장을 보여주고 낭독하게 하였다(<그림 5>). 본격적인 녹음 전 연습용으

- 10) KSB5, KSB6, KSB7은 스페인어 강세를 전혀 지키지 못하였기 때문에 산출실험 자료 분석 시 제외하였다. 스페인어는 강세가 중요한 언어로, 마지막에서 두 번째 음절에 강세가 오는 것이 원칙이기 때문에 *tate*에서 첫 번째 음절 *ta*에 강세를 주어야 하는데 KSB는 종종 마지막 음절인 *te*에 강세를 표시하고는 하였다. 강세 유무에 따른 모음 음가의 변화가 있을 수 있기 때문에(Nadeu 2014; Cobb and Simonet 2015) 혹시 모를 음가 차이를 방지하고자 명확하게 강세를 지키지 않은 참여자들은 결과 분석에서 제외하기로 하였다.
- 11) 스페인어에서 [ke, ki] 소리를 내기 위해서 ‘que, qui’로 표기하였다. ‘ke, ki’로 표기할 수도 있었겠지만 ‘k’는 거의 외래어 표기에만 사용되므로 자연스러운 철자표기를 위해 ‘q’를 선택하였다.



〈그림 5〉 프레젠테이션 슬라이드 예시

로 유사한 문장을 두 개 제시하였다. 참여자로 하여금 본인의 낭독 속도에 맞춰 직접 슬라이드를 전환할 수 있도록 하여 자연스러운 발화가 이루어지도록 하였다. Smart Recorder 어플리케이션을 사용하여 샘플링 주파수 22,050Hz로 녹음하였다.¹²⁾ 낭독 실험 전체를 수행하는 데 대략 5분 내외가 소요되었다.

34
35

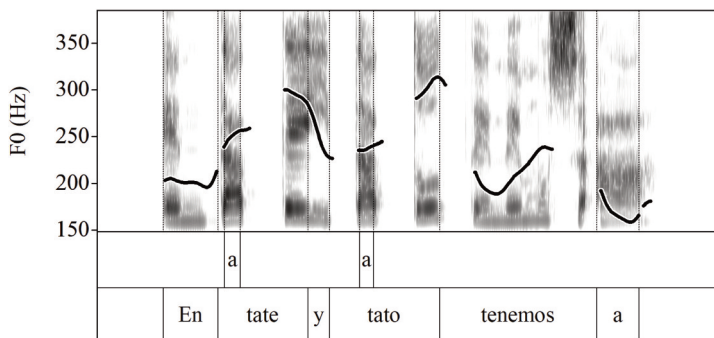
IV. 산출실험 결과 및 분석

1. 산출실험 결과 값 정리

총 24명의 참여자가 45문장씩 발화하였고, 각 문장에서 2개의 모음을 측정하였으므로 총 2,160개의 데이터가 본 논문의 분석 대상이 되었다. 녹음 자료 분석은 음향분석 프로그램인 프랏을 활용하였다. 우선, 프랏의 EasyAlign (Goldman 2011) 플러그인 프로그램을 이용하여 각 녹음파일에 대해 실험 문장별로 세그멘테이션(Macro-segmentation)된 TextGrid 파일을 만들었다. 이 TextGrid에서 본 연구의 분석 대상인 강세 있는 모음 구간을 표기하였는데, 이때 모음 구간은 전후 자음의 전이 구간을 제외한 안정된 구간으로 설정하였다. 즉, 문장들 “En C1VC2e y C1VC2o tenemos V.”에서 첫 번째 V와 두 번째 V에 해당하는 구간을 표시하되, C1, C2의 영향이 최대한 없도록 하였다.

<그림 6>은 KSA1의 녹음파일에서 임의로 선택한 하나의 문장에 대한

12) Podesva and Zsiga(2013)에 따르면 22kHz의 샘플링 레이트(sampling rate)면 대부분의 언어학 실험 연구에서 충분한 수치다.



〈그림 6〉 KSA10이 발화한 문장 “En tate y tato tenemos a.”의 TextGrid

TextGrid다. 그림과 같이 첫 번째 구간 티어에 표기한 모음 구간의 가운데 지점(50%)에서 포먼트(F1, F2, F3)를 측정하였다. 보다 효율적인 측정을 위해 스크립트(McCloy 2012)로 자동 산출하였다. 이 스크립트는 사용자가 음성 파일과 해당 TextGrid 파일이 저장되어 있는 위치를 지정하면 원하는 구간 티어의 모든 구간(interval)에 대한 50% 지점 포먼트를 자동으로 측정해 준다.¹³⁾ 이 과정 후에도 이상 치로 나타나거나 포먼트가 잘 나타나지 않는 소리는 결측값으로 처리하였다.

이렇게 측정한 포먼트는 한 번 더 후처리를 해주어 성별에 따른 차이로 발생한 주파수를 없애야 한다. 대개 남성이 긴 성도를 가지기 때문에 기본주파수가 여성에 비해 낮게 나타나기 때문이다. 이를 조정하기 위해 정규화 작업을 진행하였다(Yang 1996; 김주연 2014). 우선, F1이 600Hz를 상회하는 모음은 스페인어에서 /a/가 유일하기 때문에 /a/의 F3값을 기준으로 한다. 본 연구에서는 여성의 측정값을 정규화의 기준으로 삼았다. 여성 참여자가 비교적 많은데다 원어민 화자가 모두 여성이기 때문에 여성 측정값을 기준으로 하였다. 그리고

13) 이때 여성 자료는 5500Hz를, 남성 자료는 5000Hz를 최대 포먼트 수치로, 포먼트 수(number of formants)는 5로, 창 길이(window length)는 0.025로 설정하는 것을 기본으로 하였다. 모든 파일에 대한 자동 산출이 끝난 후 F1과 F2 이상 치(outlier)는 수작업으로 확인하였다.

선행연구들과 같이 다음 공식을 사용하였다.

$$k = \frac{\text{여성 참여자의 F3 평균값}}{\text{남성 참여자의 F3 평균값}}$$

/a/의 F3에 대한 여성의 평균값을 남성의 평균값으로 나눈 값을 k 라고 하였다.¹⁴⁾ 이때 여성 참여자의 F3 평균값은 2768.19, 남성은 2418.01로, k 를 계산하면 1.1448이 된다. 이 k 값을 남성 피험자의 모든 F1, F2 값에 곱하였다. 실험 결과 분석은 이렇게 정규화를 마친 값으로 진행하였다.

2. 산출실험 결과

산출실험 결과 분석에는 SPSS Statistics 25 프로그램(IBM Corp 2017)을 사용하였다. 우선 SPSS의 Multiple Imputation 기능을 활용하여 결측값을 적당한 수치로 대체하여 완전한 2,160개 데이터셋을 구축한 후 반복측정 다변량분산분석(Repeated-Measures Multivariate ANOVA (MANOVA))을 시행하였다.¹⁵⁾ 그리고 세 그룹의 피험자가 5개 단모음에 대해 각각 18개의 토큰(token)을 산출하였으므로 스페인어 수준(혹은 경험 정도)을 개체 간 요인(between subject factor)으로, 발화 횟수를 개체 내 요인(within subject factor)으로 설정하였다. 또한 보다 구체적인 결과를 위해 F1과 F2 각각의 사후 검정(Post-hoc Test) 결과도 제시하였다. 본고에서는 SPSS가 제시하는 결과값 중 가장 보수적인 Pillai의 궤적과 가장 널리 사용되는 Wilks의 랏다, 두 가지 결과값을 사용하였고, 사후 검정은 Bonferroni 수치를 제시하였다. 결과 분석의 편의성을 위해 실험 결과는 앞서 분류한 유사도에 따라 정리하였다.

- 14) F3값은 이상치 확인 작업을 거치지 않은 대신 전체 평균값 ± 2 *표준편차 범위에 포함되지 않는 값은 제외하고 평균값을 산출하였다. 여성 참여자 자료에서는 14개가, 남성 참여자 자료에서는 10개가 제외되었다.
- 15) 다변량분산분석은 종속변인(dependent variable)이 두 개 이상 있으며, 두 변인 간 상관관계가 존재할 때 사용한다. 본 연구는 F1과 F2 값을 분석대상으로 하며, 두 값은 독립적이지 않으므로 단순한 ANOVA보다 MANOVA가 보다 정확한 집단 간 차이를 밝히는데 유용할 것이라고 판단하였다.

1) 유사4단계 모음

유사4단계 모음은 한국어에서 그대로 전이하더라도 원어민과 학습자 간 발음 차이가 없거나 매우 적을 것이라고 예상하였다. /e/에 대한 MANOVA 분석 결과 Pillai의 꺾적=1.838 ($p=0.140$), Wilks의 람다=1.938 ($p=0.124$)로 그룹 간 유의미한 차이가 없는 것으로 나왔다(<표 5>). F1과 F2를 따로 살펴봐도 유의미한 차이가 없었다(F1: $F(2)=2.207, p>0.05$, F2: $F(2)=1.335, p>0.05$). 다만 그룹별 F1, F2 평균에서 고급 학습자 그룹이 초급에 비해 NS와 조금 더 가까운 수치를 보인다는 점은 주목할 만하다. 그럼에도 역시 Bonferroni 사후검정에서 그룹 간 대응 차의 신뢰구간을 살펴보면 모두 0을 포함하고 있어 어떤 두 그룹도 통계적으로 유의미한 차이를 가진다고 볼 수 없다.¹⁶⁾

<표 5> 그룹 간 /e/ 포먼트 값에 대한 MANOVA 결과

	Value	F	df	p	η_p^2
Pillai's Trace	.298	1.838	4, 42	.140	.149
Wilks' Lambda	.702	1.933b	4, 40	.124	.162

<표 6> 그룹별 /e/ F1 값에 대한 Bonferroni 사후검정

대 응	대응차(평균, 표준오차)	p	신뢰구간(95%)
NS-KSA	-11.786 (21.972)	1.000	(-68.944, 45.372)
NS-KSB	-42.313 (23.305)	.251	(-102.939, 18.312)
KSA-KSB	-30.528 (17.371)	.280	(-75.715, 14.659)

<표 7> 그룹별 /e/ F2 값에 대한 Bonferroni 사후검정

대 응	대응차(평균, 표준오차)	p	신뢰구간(95%)
NS-KSA	50.327 (88.936)	1.000	(-181.026, 281.681)
NS-KSB	139.982 (94.331)	.458	(-105.405, 385.370)
KSA-KSB	89.655 (70.310)	.649	(-93.246, 272.556)

16) 수치는 소수점 세 자리까지 제시하였다.

2) 유사3단계 모음

매우 유사한 유사3단계 모음은 한국어에서 그대로 전이 시 NS와 비교했을 때 약간의 차이가 있을 수 있지만. 그래도 상당히 유사할 것으로 예상된다. 우선 /a/ 포먼트 값에 대한 MANOVA 결과부터 알아본다.¹⁷⁾ 두 통계 값 모두 $p>0.05$ (Pillai의 궤적=0.384 ($p=0.068$), Wilks의 람다=0.633 ($p=0.064$))가 되어 통계적으로 유의미한 차이가 없는 것으로 나왔다(<표 8>). 하지만 그룹별 F1, F2 값에 대한 개체 간 요인 효과 결과에서는 그룹 간 F2값의 유의미한 차이가 있었다(F1: F(2)=0.754, $p>0.05$, F2: F(2)=4.955, $p=0.018$). 그리고 이 차

〈표 8〉 그룹 간 /a/ 포먼트 값에 대한 MANOVA 결과

	Value	F	df	<i>p</i>	η_p^2
Pillai's Trace	.384	2.377	4, 40	.068	.192
Wilks' Lambda	.633	2.437	4, 38	.064	.204

〈표 9〉 그룹별 /a/ F1 값에 대한 Bonferroni 사후검정

대응	대응차(평균, 표준오차)	<i>p</i>	신뢰구간(95%)
NS-KSA	12.487 (47.999)	1.000	(-112.914, 137.887)
NS-KSB	52.448 (50.341)	.930	(-79.074, 183.969)
KSA-KSB	39.961 (33.940)	.759	(-75.715, 14.659)

〈표 10〉 그룹별 /a/ F2 값에 대한 Bonferroni 사후검정

대응	대응차(평균, 표준오차)	<i>p</i>	신뢰구간(95%)
NS-KSA	160.532 (58.701)	.038	(7.170, 313.894)
NS-KSB	190.516 (61.566)	.017	(29.668, 351.364)
KSA-KSB	29.984 (41.508)	1.000	(-78.460, 138.427)

17) 처음 MANOVA를 실시했을 때 Pillai의 궤적=0.474 ($p=0.020$), Wilks의 람다=0.534 ($p=0.012$)로 예상과 달리 세 그룹 간 포먼트 값의 유의미한 차이가 있었다. 그런데 NS3의 값이 다른 NS 화자들과 크게 다르다는 점을 발견하였다(NS1의 /a/ 포먼트 측정값 평균은 (F1, F2)=(826.09, 1712.67), NS2는 (F1, F2)=(716.83, 1669.98), NS4는 (F1, F2)=(864.97, 1634.45)인 반면에 NS3은 (F1, F2)=(1100.78, 1811.00)). 이에 따라 NS3을 제외한 세 명의 데이터로 MANOVA를 재실시하였고, 본고에서 제시한 수치들은 이에 대한 결과다.

이는 NS와 KSA 사이($p=0.38$)보다 NS와 KSB 사이($p=0.017$)에 더 크게 나타났다. 즉, 유사4단계에서 보다 확실한 발음 차이가 존재하는 것으로 나타났으며, KSB에 비해 KSA가 NS와 조금 더 가까운 수치를 보였다. 하지만 동시에 KSA와 KSB 간 차이는 없는 것으로 나타나($p=1.000$) 학습 효과를 뒷받침하는 확실한 근거는 될 수 없다.

다음으로 <표 11>는 /i/ 포먼트 값에 대한 MANOVA 결과를 정리한 것이다. /i/도 역시 세 그룹 간 포먼트 평면도 상 유의미한 차이는 없는 것으로 나타났다. F1, F2값을 별도의 변수로 계산했을 때도 모두 $p>0.05$ (F1: $F(2)=2.431$, $p>0.05$, F2: $F(2)=1.130$, $p>0.05$)로 나타나 원어민과 학습자 간 /i/ 발음은 거의 동일하다고 볼 수 있겠다. <표 12>, <표 13>에 나와 있는 그룹 간 분석에서도 p 값이 워낙 크게 나타나 사실상 세 그룹 간 /i/의 높이(height)나 전설 정도(frontness)나 차이는 없다고 봐야 한다.

<표 11> 그룹 간 /i/ 포먼트 값에 대한 MANOVA 결과

	Value	F	df	p	η_p^2
Pillai's Trace	.315	1.960	4, 42	.118	.157
Wilks' Lambda	.706	1.905	4, 40	.128	.160

<표 12> 그룹별 /i/ F1 값에 대한 Bonferroni 사후검정

대응	대응차(평균, 표준오차)	p	신뢰구간(95%)
NS-KSA	-48.013 (23.971)	.175	(-110.370, 14.343)
NS-KSB	-19.440 (25.425)	1.000	(-85.579, 46.670)
KSA-KSB	28.574 (18.951)	.440	(-20.724, 77.871)

<표 13> 그룹별 /i/ F2 값에 대한 Bonferroni 사후검정

대응	대응차(평균, 표준오차)	p	신뢰구간(95%)
NS-KSA	70.757 (112.188)	1.000	(-221.084, 362.598)
NS-KSB	167.890 (118.994)	.519	(-141.654, 477.434)
KSA-KSB	97.133 (88.693)	.858	(-133.588, 327.854)

3) 유사2단계 모음

다음은 유사2단계인 /u/에 대한 MANOVA 결과를 살펴본다. 비록 유사2단계로 분류하였지만 /u/는 사실 한국어 /ㅜ/와 절대적인 포먼트 평면도 거리상 가까운 편이기 때문에 한국어를 그대로 전이시키더라도 NS 발음과 큰 차이가 없을 수 있어 산출 면에 있어서는 오류가 발생하지 않을 수도 있다. 실제로 Pillai의 꺾적=0.297 ($p=0.141$), Wilks의 람다=0.716 ($p=0.144$)으로 그룹 간 차이가 없었다(<표 14>). F1과 F2를 분리해서 분석하더라도 $p>0.05$ 였으며, 게다가 모든 비교 쌍에서 $p>0.05$ 로 나왔기 때문에 그룹 간 유의미한 차이가 있다고 볼 수 없었다.

〈표 14〉 그룹 간 /u/ 포먼트 값에 대한 MANOVA 결과

	Value	F	df	p	η_p^2
Pillai's Trace	.297	1.829	4, 42	.141	.148
Wilks' Lambda	.716	1.819	4, 40	.144	.154

〈표 15〉 그룹별 /u/ F1 값에 대한 Bonferroni 사후검정

대 응	대응차(평균, 표준오차)	p	신뢰구간(95%)
NS-KSA	-57.891 (25.979)	.111	(-125.473, 9.691)
NS-KSB	-46.124 (27.555)	.327	(-117.805, 25.557)
KSA-KSB	11.767 (20.539)	1.000	(-41.661, 65.195)

〈표 16〉 그룹별 /u/ F2 값에 대한 Bonferroni 사후검정

대 응	대응차(평균, 표준오차)	p	신뢰구간(95%)
NS-KSA	-86.870 (80.479)	.878	(-296.226, 122.484)
NS-KSB	-144.408 (85.361)	.316	(-366.463, 77.646)
KSA-KSB	-57.537 (63.624)	1.000	(-223.047, 107.972)

4) 유사1단계 모음

마지막으로 한국어-스페인어 간 가장 큰 차이를 보인 모음인 /o/에 대한 산출실험 결과를 제시한다. <표 17>에 따르면, Wilks의 람다=0.577 ($p=0.024$)

는 물론 비교적 보수적인 Pillai의 꺾적=0.424 ($p=0.037$)도 세 그룹 간 포먼트 값의 유의미한 차이가 있을 거라는 예상에 부합했다. 이를 보다 구체적으로 살펴보기 위해 F1과 F2 각 값에 대한 개체 간 요인 효과를 계산했을 때도 역시 유의미한 p 값이 도출되었다(F1: $F(2)=5.503, p=0.012$, F2: $F(2)=4.697, p=0.021$). NS의 F1, F2 모두 학습자 포먼트 값보다 더 크게 나타났다. 즉, NS의 /o/가 비교적 아래, 그리고 앞쪽에서 발음된다고 볼 수 있다. 이는 역시 /o/ 보다 입천장 뒤쪽에서 발음되는 /ɯ/ 한국어 모음의 영향으로 볼 수 있다.

또한, <표 18>, <표 19>에 정리된 사후검정 결과에 따르면 통계적으로 유의미한 차이는 NS와 KSA, NS와 KSB의 F1 값 간에, 그리고 NS와 KSB의 F2 값 간에 발견됐다. 추가적으로 NS와 KSA의 F2 값도 실질적으로 차이가 있을 확률이 높다고 할 수 있다($p=0.054$). 비록 KSA와 KSB 간 F1, F2 값이 차이가 전혀 없는 것으로 나타났지만($p=1.000$) 그룹별 평균이나 p 값을 고려할 때 KSA가 비교적 원어민 화자에 근접하게 발음하는 것으로 보인다(F1: $0.025>0.013$, F2: $0.054>0.021$). 이러한 결과는 학습 효과를 뒷받침하는 증거로 해석할 수 있다.

<표 17> 그룹 간 /o/ 포먼트 값에 대한 MANOVA 결과

	Value	F	df	p	η^2
Pillai's Trace	.424	2.825	4, 42	.037	.212
Wilks' Lambda	.577	3.166	4, 40	.024	.240

<표 18> 그룹별 /o/ F1 값에 대한 Bonferroni 사후검정

대 응	대응차(평균, 표준오차)	p	신뢰구간(95%)
NS-KSA	62.685 (21.522)	.025	(6.698, 118.672)
NS-KSB	72.650 (22.828)	.013	(13.267, 132.033)
KSA-KSB	9.965 (17.015)	1.000	(-34.296, 54.227)

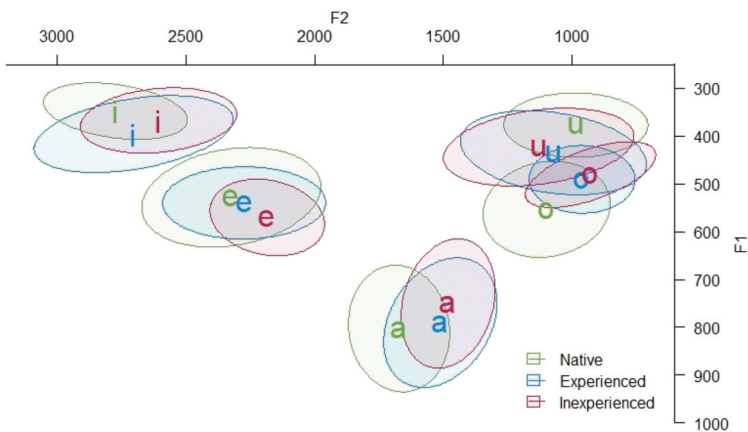
<표 19> 그룹별 /o/ F2 값에 대한 Bonferroni 사후검정

대 응	대응차(평균, 표준오차)	p	신뢰구간(95%)
NS-KSA	137.655 (53.593)	.054	(-1.759, 277.069)
NS-KSB	170.453 (56.844)	.021	(22.582, 318.324)
KSA-KSB	32.798 (42.369)	1.000	(-77.419, 143.015)

5) 산출실험 결과 분석

다섯 모음의 산출실험 결과를 요약하면서 이번 장을 정리한다. <그림 7>은 R 프로그램의 phonR 패키지(McCloy 2016)를 활용하여 그룹별 산출실험 결과를 포먼트 평면도에 나타낸 것이다. 초록색은 NS, 파란색은 KSA, 빨간색은 KSB의 값을 나타내며, F1과 F2의 평균값을 중심으로 $\pm 1SD$ (표준편차)가 포함되는 범위를 그린 것이다(즉 약 68%의 신뢰타원(Confidence Ellipse)이다). 개별 타원의 크기가 클수록 해당 그룹의 모음 발음이 피험자 간 다양하게 나타난다고 볼 수 있으며, 타원 간 넓이가 많이 겹칠수록 서로 유사하게 발음한다고 해석할 수 있다. 즉, /i/, /u/의 경우 원어민에 비해 학습자가 다양하게 발음하는 반면 /e/는 원어민이 비교적 일관성 없는 발음을 구사했다고 할 수 있다. 그리고 /e/, /a/의 경우 초록 원이 빨간 원보다 파란 원과 겹치는 범위가 넓으므로 고급 학습자가 원어민과 비교적 유사하게 발음했다고 분석할 수 있다. 특히 /e/의 경우 KSA와 NS의 타원이 거의 포개어짐을 확인할 수 있다.

한편, /o/는 학습자 두 그룹의 타원끼리는 상당부분 겹치는데 반해 원어민 화자는 한참 아래에 위치한다. /u/의 경우 KS의 타원들이 NS의 타원을 일부 포함하고 있기는 하지만 학습자 간 겹치는 부분이 훨씬 크게 나타났다. 그리고



〈그림 7〉 산출실험 그룹별 결과

/o/와 /u/를 비교했을 때 NS의 두 모음 범위는 겹치지 않는 반면에 학습 경험량과 관계없이 KS의 두 모음은 겹쳐져 나타난다. 이는 한국인 학습자가 /u/는 원어민보다 아래-앞에서 발음하는 동시에 /o/는 보다 후설 고모음으로 발음하기 때문이다. 이렇다보니 학습자의 /u/는 원어민의 /o/ 범위와 겹치게 되고, 또 /o/는 원어민의 /u/ 범위와도 겹쳐있게 된다는 문제가 생긴다. 즉, 학습자의 /o/가 /u/로 인식될 수도, 반대로 /u/가 /o/로 인식될 수도 있다는 것이다. 그리고 이러한 분석은 앞의 장에서 살펴본 인식실험 결과와도 어느 정도 같은 맥락이다.

V. 논 의

1. L1 유사도에 따른 발음 정확도

앞서 발음실험을 통해 한국어 모음과의 유사도에 비례하여 한국인 학습자의 스페인어 모음 발음이 원어민 수준과 유사해지지는지 알아보았다. 한국인 모어 화자에게 스페인어 모음 중 ‘같은 소리’에 가까운 /e/, /a/, /i/는 그대로 한국어에서 전이시키더라도 산출 측면에서 거의 문제가 없을 것으로, 그 외 /u/와 /o/는 원어민과 차이를 보일 것으로 예상했다.

산출실험의 MANOVA 결과도 어느 정도 유사 단계와 관련성을 보였다. 유사 4단계 /e/가 상당히 높은 p 값을 기록하여 원어민과 학습자 간 발음 상 거의 차이가 없었으며, 유사1단계 /o/는 가장 낮은 p 값으로 원어민과 학습자 간 유의미한 차이가 존재하였다. 한편, 유사3단계인 /a/, /i/도 원어민과 학습자 간 통계적으로 유의미한 발음상의 차이가 없는 것으로 나타났다. 다만 /a/의 경우 0.05에 가까운 값이 도출됐는데, 이러한 차이는 F1보다는 F2에서 비롯된 것으로 보인다(F2 값에 대해 그룹 간 유의미한 차이가 있는 것으로 나타남($p=0.018$)). 원어민에 비해 학습자가 /a/를 훨씬 더 입 뒤쪽에서 발음했다는 것이다.¹⁸⁾

18) 그런데 이는 사실 앞서 제시한 한국어-스페인어 모음평면도의 내용과는 차이가 있는 결과로, 후속 연구로 같은 기준 하에서 보다 정확하게 측정한 한국어와 스페인어 모음 위치를 알아볼 필요가 있다.

〈표 20〉 한국어-스페인어 모음 간 거리와 스페인어 모음 포먼트 값에 대한 MANOVA 결과

	a	e	i	o	u
한국어 대응 모음과의 거리(남성)	131.33	25.80	145.96	246.24	100.37
한국어 대응 모음과의 거리(여성)	191.56	161.23	202.43	448.71	162.94
<i>p-value</i> (Pillai's Trace)	.068	.140	.118	.037	.141
<i>p-value</i> (Wilks' Lambda)	.064	.124	.128	.024	.144

그리고 /u/의 *p*값이 의외로 /e/에 상응하는 높은 값이었는데, 이를 이해하기 위해 앞서 살펴본 유클리드 거리를 활용하였다. 비록 /u/는 한국어 모음 두 개(/ㅜ/, /ㅡ/)와 가까운 조음위치를 가지기 때문에 한국인 학습자의 인식 면에서 혼란을 일으킬 가능성이 다분하지만, 그와 동시에 대응 모음 /ㅡ/와 매우 가깝기 때문에 산출실험에서는 의외로 원어민과 학습자 간 차이가 없다는 결과가 나온 것으로 보인다. 요약하자면 스페인어 모음을 발음하는 데 있어 한국인 학습자는 한국어 모음을 일대일 대응시키기 때문에 한국어 모음과 포먼트 평면도 상 가까이 위치한 스페인어 모음일수록 원어민과 유사하게 발음하는 경향성을 보인다.

2. L2 경험량에 따른 모음 발음 정확도

다음으로, 스페인어 경험량에 따라 단모음 발음이 개선되는지, 그리고 그 개선 정도가 모국어 화자 수준에 이르게 되는지에 대해 논한다. 이 둘은 피험자 세 그룹을 비교하여 분석하면 답할 수 있는 문제다.

산출실험 결과에 따르면 학습효과를 지지하는 실질적인 증거를 찾을 수 없었다. 스페인어 모음별로 F1, F2에 대한 개체 간 요인 효과의 *p*값과 각 그룹 수치의 평균을 통해 세 그룹을 비교해 보았다(<표 21>). 우선 유사4단계 모음 /e/의 F1, F2는 그룹 간 통계적으로 유의미한 차이가 없는 것으로 나타났다. 그래도 그룹별 평균을 비교해보면 KSA 그룹의 수치가 KSB에 비해 NS에 가까웠다. 역시 /i/의 F1, F2에 대해서도 그룹 간 차이가 없는 것으로 나타났다. 그런데 KSB의 F1 평균값은 KSA에 비해 NS에 더 가까웠다. 한편, 또 다른 유

〈표 21〉 그룹별 스페인어 모음 F1, F2 평균값 비교

스페인어 모음		NS	KSA	KSB	그룹별 비교
a	F1	802.630	790.144	750.183	NS>KSA>KSB
	F2	1672.365	1511.833	1481.849	NS>KSA>KSB*
e	F1	527.849	539.635	570.163	NS<KSA<KSB
	F2	2324.685	2274.357	2184.702	NS>KSA>KSB
i	F1	347.609	395.622	367.048	NS<KSB<KSA
	F2	2774.890	2704.133	2607.000	NS>KSA>KSB
o	F1	553.065	490.380	480.415	NS>KSA>KSB*
	F2	1099.147	961.492	928.694	NS>KSA>KSB*
u	F1	376.762	434.654	422.886	NS<KSB<KSA
	F2	985.941	1072.812	1130.349	NS<KSA<KSB

* $p<0.05$

사3단계 모음 /a/의 F1은 그룹 간 차이가 없는 반면 F2는 유의미한 차이가 있었다($p=0.018$). 특히 학습자 그룹끼리는 차이가 없는 것으로 나타났지만 ($p=1.000$) 학습자와 원어민 간에는 유의미한 차이가 존재했다(NS-KSA: $p=0.038$; NS-KSB: $p=0.017$). 이는 학습자의 경험량 증가에 따른 발음 개선 효과가 없었음을 보여주는 결과로, 학습효과를 오히려 반증한다고 볼 수 있다.

그리고 유사2단계 모음 /u/는 MANOVA 결과에서와 마찬가지로 F1, F2값에 대해서도 그룹별 유의미한 차이를 보이지 않았다. 그런데 F1, F2값에 있어서 두 학습자 그룹에 대한 p 값이 1.000에 달한다는 점을 고려할 때 학습량 증가에 따른 스페인어 발음 변화가 없었다고 분석할 수 있겠다.

마지막으로 유사1단계 모음 /o/ 결과를 알아본다. MANOVA 결과부터 통계적으로 유의미하게 나왔던 /o/는 역시 F1, F2값 모두에 대해 그룹 간 차이가 있는 것으로 드러났다. 비록 KSB 그룹보다는 KSA가 NS와 가깝게 나타났기 때문에 학습효과를 무시할 수 없지만, 그룹 간 비교한 p 값에 따르면 유의미한 차이는 학습자와 원어민 간에만 존재했다. 요약하자면, 학습효과를 뒷받침하는 유효한 증거를 위해서는 NS와 KSB 간 차이가 존재함과 동시에 NS와 KSA 간 차이가 없거나 적어도 세 그룹 간 유의미한 차이가 있어야 한다. 하지만 이번 실험 결과에 따르면 세 그룹 모두 차이가 없다고 나오거나 원어민과 학

습자 간 차이만 있었다. 즉, 스페인어 경험량이 5년 이상인 학습자라도 원어민의 모음 인식/발음 수준에 도달하지 못했으며, 사실상 경험량이 1년 미만인 학습자와 큰 차이를 보이지 않았다.

VI. 결론 및 한계점

본 연구는 지금까지 간과되어 온 한국인 스페인어 모음 발음 양상을 음성학 실험을 통해 알아보았다. 우선, 한국어와 스페인어 모음체계를 비교하여 그 차이를 발견하였다. 이를 토대로 간단한 청취실험을 통해 실제로 스페인어 모어 화자가 한국인 학습자의 /o/를 /u/로 잘못 인식하는 현상을 확인하였다. 그리고 스페인어 모어 화자와 한국인 학습자들이 발화한 스페인어를 음향 분석 프로그램을 통해 과학적으로 분석하였다. 그 결과, 한국인 학습자들은 스페인어 수준과 관계없이 스페인어 모어 화자의 모음 발음과 유의미한 차이를 보였다. 이러한 차이는 한국어와 스페인어가 서로 다른 모음체계를 가지고 있음에도 스페인어를 처음 학습 시 두 모음 체계를 일대일 대응시켜 발음 연습하는데서 비롯된다. 게다가 어느 정도 학습이 이루어진 후에도 그 차이를 스스로 인지할 수 있기가 쉽지 않을뿐더러 이를 지적해주는 교수자가 없기 때문에 학습량이 증가함에도 발음이 개선되지 않는다. 따라서 본 연구의 일차적인 의의는 한국인 스페인어 학습자 및 교수자가 모음 인식 및 발음 상 오류를 인지하게끔 하는데 있다. 그리고 실험을 통해 수치적으로 그 차이를 제시하였기 때문에 이를 토대로 발음 개선 방향을 구상하기 용이하다.

하지만 한국인의 스페인어 모음 습득 분야에서 처음으로 실험음성학적 방법론을 도입한 연구인만큼 다수의 한계점을 가진다. 우선, ‘/p, t, k/+V’의 제한된 환경에서만 산출 실험을 했는데 추후엔 보다 다양한 환경에서의 실험이 필요하다. 그리고 본 연구에서는 편의를 위해 강세가 있는 모음만 실험 대상에 포함하였는데, 한국인 학습자들이 잘 지키지 못하는 스페인어 강세의 유무에 따라 스페인어 원어민과 학습자의 모음 발음이 어떠한 양상을 보이는지 알아

보는 것도 중요한 연구 과제라고 할 수 있겠다.

또한, 추후에는 L1 한국어의 영향을 보다 명확하게 규명하는 것도 필요하다. 이를 위해 일차적으로 학습자들에게 한국어 실험도 동일하게 진행, 데이터를 수집하여 분석해 볼 수 있겠다. 한국인 학습자들이 발음한 한국어 모음과 스페인어 모음을 비교해보는 것이다. 그리고 더 나아가 스페인어 학습 경험이 전무한 한국어 화자의 데이터와 비교를 위해 한국어 인식 및 산출 실험을 추가하여 학습자들의 L2 학습으로 인한 L1 발음의 변화를 알아볼 수도 있을 것이다. 이처럼 한국인 스페인어 학습자의 모음 습득 분야와 관련, 다양한 추가 연구가 기대된다.

참고문헌

- 강현화·신자영·이재성·임효상(2003), 『대조분석론: 한국어·스페인어 문형 대조를 바탕으로』, 역락.
- 권성미(2009), 『한국어 발음 습득 연구: 모음 중심의 실험음성학적 연구』, 박이정.
- 김승한(2009), 『국내 스페인어 학습자들의 발음 오류에 관한 연구』, 석사학위논문: 조선대학교.
- 김원필(2001), 「스페인어 자음의 발음 형성기술 연구」, 서어서문연구, Vol. 20, pp. 3-24.
- (2009), 「스페인어 모어 화자의 한국어 분절음 학습방안: 양 언어의 대조분석을 중심으로」, 이베로아메리카연구, Vol. 20, No. 1, pp. 233-260.
- (2013), 「스페인어 교육학: 한국어 간섭으로 인한 스페인어 발음 오류현상」, 스페인어문학, Vol. 67, pp. 177-197.
- (2014), 「스페인어 음성학 관련 용어 및 개념의 혼동현상: 국내 스페인어 학습자들의 경우를 중심으로」, 스페인라틴아메리카연구, Vol. 7, No. 1, pp. 107-129.
- (2017), 「스페인어 변이음의 효율적 교수 방안: 자음을 중심으로」, 스페인어문학, Vol. 84, pp. 35-54.
- 김주연(2014), 「중국인 학습자의 한국어 모음 습득에 대한 제2언어 습득 모델 비교 연구」, 말소리와 음성과학, Vol. 6, No. 4, pp. 27-36.

- 문정(2014), 「한국인 스페인어 학습자의 발화에 대한 음향음성학적 연구: /p, t, k/ 및 /b, d, g/를 중심으로」, 석사학위 논문, 고려대학교.
- 서경석(2007), 「스페인어권 학습자에 대한 한국어 발음 교육 방안」, 이베로아메리카, Vol. 9, No. 2, pp. 197-220.
- 윤영도(2007), 「합성한 영어모음에 대한 한국인들과 영어원어민들의 인식 비교」, 현대영어영문학, Vol. 51, No. 2, pp. 117-132.
- 이신영(2015), 「중남미 스페인어권 학습자의 한국어 단모음 발음 연구」, 석사학위 논문, 경희대학교.
- Boersma, Paul and David Weenink(2018), *Praat: doing phonetics by computer* [Computer program], Version 6.0.37, <http://www.praat.org>.
- Bohn, Ocke-Schwen and James Emil Flege(1990), “Interlingual identification and the role of foreign language experience in L2 vowel perception”, *Applied Psycholinguistics*, Vol. 11, No. 3, pp. 303-328.
- _____(1992), “The Production of New and Similar Vowels by Adult German Learners of English”, *Studies in Second Language Acquisition*, Vol. 14, pp. 131-158.
- Catford, John Cunnison(2001), *A Practical Introduction to Phonetics*, NY: Oxford University Press.
- Clegg, J. Halvor and Willis C. Fails(2017), *Manual de fonética y fonología españolas*, Routledge.
- Chládková, Kateřina, Paola Escudero and Paul Boersma(2011), “Context-specific acoustic differences between Peruvian and Iberian Spanish vowels”, *Journal of the Acoustical Society of America*, Vol. 130, No. 1, pp. 416-428.
- Cobb, Katherine and Miguel Simonet(2015), “Adult Second Language Learning of Spanish Vowels”, *Hispania*, Vol. 98, No. 1, pp. 47-60.
- Flege, James Emil(1987), “The production of “new” and “similar” phones in a foreign language: Evidence for the effect of equivalence classification”, *Journal of Phonetics*, Vol. 15, pp. 47-65.
- _____(1992), “Speech learning in a second language”, Charles A. Ferguson, Lise Menn, & Carol Stoel-Gammon(Eds.), *Phonological Development: Models, Research, and Application*, Timonium, MD: York Press, pp. 565-604.

- _____(1995), “Second language speech learning theory, findings, and problems”, Winifred Strange(Ed.), *Speech Perception and Linguistic Experience: Issues in cross-language research*, Timonium, MD: York Press, pp. 233-277.
- Flege, James Emil, Ocke-Schwen Bohn and Sunyoung Jang(1997), “Effects of experience on non-native speakers’ production and perception of English vowels”, *Journal of Phonetics*, Vol. 25, No. 4, pp. 437-470.
- Flege, James Emil and James Hillenbrand(1984), “Limits on phonetic accuracy in foreign language speech production”, *The Journal of the Acoustical Society of America*, Vol. 76, No. 3, pp. 708-721.
- Goldman, Jean-Philippe(2011), “EasyAlign: an automatic phonetic alignment tool under Praat”, *Proceedings of InterSpeech*, September 2011.
- Hualde, José Ignacio(2005), *The Sounds of Spanish*, Cambridge: Cambridge University.
- IBM Corp.(2017), *IBM SPSS Statistics for Windows*, Version 25.0, NY: IBM Corp.
- Lee, Sue Ann S. and Gregory K. Iverson(2012), “Vowel Category Formation in Korean-English Bilingual Children, Journal of Speech”, *Language, and Hearing Research*, 55, pp. 1449-1462.
- McCloy, Daniel(2012), PRAAT SCRIPT “SEMI-AUTO FORMANT EXTRACTOR”, <https://github.com/drammock/praat-semiauto/blob/master/SemiAutoFormantExtractor.praat>.
- _____(2016), *phonR: tools for phoneticians and phonologists*, R package version 1.0-7.
- Morrison Geoffrey Stewart(2003), “Perception and Production of Spanish Vowels by English Speakers”, Maira-Josep Solé, Daniel Recansens, and Joaquin Romero(Eds.), *Proceedings of the 15th International Congress of Phonetic Sciences*, Barcelona, pp. 1533–1536.
- Nadeu, Marianna(2014), “Stress- and speech rate-induced vowel quality variation in Catalan and Spanish”, *Journal of Phonetics*, Vol. 46, pp. 1-22.
- Podesva, Robert J. and Elizabeth Zsiga(2013), “Sound recordings: acoustic and articulatory data”, Robert J. Podesva and Devyani Sharma(Eds.), *Research Methods in Linguistics*, NY: Cambridge University, pp. 169-194.
- Shin, Jiyoung, Jieun Kiaer and Jaeun Cha(2012), *The Sounds of Korean*,

Cambridge: Cambridge University Press.

Yang, Byunggon(1996), “A comparative study of American English and Korean vowels produced by male and female speakers”, Journal of Phonetics, Vol. 24, pp. 245-261.

김 주 경

서울특별시 관악구 관악로1 서울대학교 3동 429호
anita92@snu.ac.kr

논문투고일: 2020년 3월 22일

심사완료일: 2020년 4월 17일

게재확정일: 2020년 4월 17일

Acoustical Analysis of Spanish Vowel Production by Korean Learners

Joo Kyeong Kim

Seoul National University

Kim, Joo Kyeong(2020), "Acoustical Analysis of Spanish Vowel Production by Korean Learners", *Revista Asiática de Estudios Iberoamericanos*, 31(1), 23-52

Abstract In this study, methods in experimental phonetics were implemented in order to compare Korean and Spanish vowel system and to scientifically analyze vowel production by Korean learners of Spanish. According to Flege's Speech Learning Model (henceforth SLM; Flege 1995), L2 sounds may be classified acoustically into three categories by their L1 counterpart availability: "identical", "similar", or "new". However, it is more appropriate to consider 'similarity' as a relative concept rather than an absolute one. Adopting Kwon (2007), Spanish vowels were classified into four levels, by the degree of similarity to Korean. We sought to investigate this L1 similarity level would work positively towards Spanish vowel acquisition by Korean. Furthermore, the amount of Spanish experience was taken into account to verify the effects of foreign language experience on vowel production.

We first noticed the possibility of miscommunication due to perception errors in Spanish vowels by Korean speakers. Then, production experiment was carried out in order to exactly locate where the Spanish vowels are produced. In total, 24 participants each read 45 Spanish sentences aloud. Praat program was utilized for formant analysis. The statistical results show that there are almost no differences in production between the three groups for vowels with relatively high similarity levels (/a/, /e/, /i/, /u/), but significant difference was found for /o/. Meanwhile, no significant difference was found between the two learner groups, even for /o/. These errors in production are possibly due Spanish instruction based on the wide-spread and deeply ingrained misconception that Spanish vowels have one-to-one correspondence with Korean.

Key words L2 Spanish, vowel acquisition, vowel production, production errors, experimental phonetics