

タイトル	The effects of pronunciation training on the development of second language phonemic categories
著者	Matthews, John
引用	北海学園大学人文論集, 9: 129-149
発行日	1997-10-31

The effects of pronunciation training on the development of second language phonemic categories

John Matthews

Abstract

This paper reports a study investigating the influence of pronunciation training on the development of second language segmental categories. Japanese learners of English were given explicit instruction in the precise articulation of the English segments [θ], [f], [v], [l] and [ɹ]. Although they were given no perceptual models in the course of training, subjects' performance on an AX discrimination task exhibited significant improvement relative to controls. This improvement, however, was not uniform across all segmental contrasts tested ([s]~[θ], [s]~[f], [b]~[v], [l]~[ɹ], [θ]~[f], [p]~[b]). It is proposed that the asymmetry in success of training is due to properties of the native language phonological system. The process of acquiring the system of segmental oppositions in the first language, which can be illustrated in theoretical terms in a feature geometry, imposes on an individual's perceptual system the specific boundaries within which categories are perceived. The phonological knowledge that is characterized by the feature geometry restricts the range of new segmental representations a learner will be able to acquire. Neither exposure to perceptual input nor explicit training in production can trigger the development of segmental representations that require elements absent from the L1 system.

Keywords: perception, pronunciation, phonemic contrasts, training

In this paper we will investigate whether pronunciation training can influence the perception of second language contrasts in an effort to address the more general question of whether second language learners can acquire novel segmental contrasts. If explicit instruction in the articulation of non-native segments contributes to the development of a new underlying phonetic (or phonological) category, then learners should be able to take advantage of this new information when discriminating perceptually members of the new category from members of other categories.

A considerable body of literature has been devoted to investigating whether Japanese learners of English can acquire the non-native [l] ~ [ɭ] contrast. The general conclusion is that Japanese and English listeners do not perceive the [l] ~ [ɭ] contrast in the same way (Miyawaki, Strange, Verbrugge, Liberman, Jenkins, and Fujimura, 1975; MacKain, Best, and Strange, 1981; Mochizuki, 1981; Mann, 1985). In addition, Japanese listeners exhibit significant sensitivity to the phonetic environment in which a segment occurs. When tested on this contrast in syllable final position, they perform much better than when they are asked to make a judgment on segments in syllable initial position or within an initial consonant cluster (Goto, 1971; Henly and Sheldon, 1986; Mochizuki, 1981; Sheldon and Strange, 1982).

Given the generally poor ability Japanese learners seem to have in acquiring the English [l] ~ [ɭ] contrast, several researchers have investigated whether it might be possible to train them to perceive the contrast. Using synthetically produced stimuli in a discrimination training protocol, Strange and Dittman (1984) found that although subjects improved in their ability to discriminate the experimental pairs they were trained on, the newly acquired ability did not generalize to other synthetic tokens or to naturally produced stimuli. Lively, Pisoni,

and Logan (1992) achieved greater success with an identification training regimen in which the contrast under study was presented in a range of phonetic contexts in a variety of tokens by different voices in natural speech. Nevertheless, these researchers report a high degree of stimulus-dependence in the performance of their subjects. Although the variability in training materials seems to have aided learners in the development of more robust categories, they continued to perform better on stimuli they had heard before as well as when stimuli were produced by a voice they had heard before.

In studying the ability to acquire non-native segmental contrasts among Japanese learners of English, it is important to recognize that English contains several non-Japanese contrasts other than the oft scrutinized [l]~[ɭ]. There are three different ways that a second language segmental contrast might correspond to the native inventory. Table 1 illustrates these correspondences with respect to L2 English contrasts and an L1 Japanese inventory. First, both members of the contrast might also exist in the native inventory. In such a case, the L2 contrast is effectively an L1 contrast that happens also to be part of the second language inventory (e.g., [p]~[b], [t]~[s]). Second, one member of an L2 contrast may be present in the L1 inventory while the

Table 1. Relationship between English contrasts and Japanese inventory

English segmental contrast	present in Japanese	absent from Japanese
[p]~[b] [t]~[s]	[p] [b] [t] [s]	
[b]~[v] [s]~[θ]	[b] [s]	[v] [θ]
[θ]~[f] [l]~[ɭ]		[θ] [f] [l] [ɭ]

other member is not (e.g., [b]~[v], [s]~[θ]). Third, both members of the L2 contrast may be absent from the L1 inventory (e.g., [θ]~[f], [l]~[ɹ]).

Three different types of evidence have been used to indicate when learners have knowledge of a segmental contrast. First, successful performance on discrimination tasks indicates that learners can recognize two acoustic cues to be members of the same category or members of two separate categories. Second, successful performance on identification tasks indicates that learners can recognize an acoustic cue to be an instance of a particular segment. Finally, accurate production of two distinct segments indicates that learners can distinguish them as members of two separate categories.

Although these sources of evidence indicate when learners have acquired knowledge of an underlying segmental contrast, unsuccessful performance on such tasks do not equally indicate a lack of knowledge of a segmental contrast. It is entirely plausible that an individual may be unable to accurately distinguish segments of two separate categories in production yet be able to discriminate them perceptually with no difficulty. Such an individual clearly has knowledge of the segmental contrast despite his or her inability to articulate them distinctly. Similarly, an individual may be unable to accurately identify which segmental category a particular acoustic cue belongs to, yet have no difficulty at recognizing two acoustic cues as being members of the same category or of different categories. This individual also demonstrates knowledge of the segmental contrast at some level. If, however, an individual is unable to recognize two acoustic cues that correspond to two separate categories as being different from one another, then there is no evidence that the individual has knowledge of the segmental contrast. Thus, unlike evidence from production tasks or

identification tasks, unsuccessful performance on discrimination tasks indicates that an individual lacks knowledge of an underlying segmental contrast.

Method

In order to test the effects of pronunciation training on the perception of non-native segmental contrasts, six English segmental contrasts were chosen for use with Japanese learners of English. Two contrasts contained members that are both absent from the Japanese segmental inventory ($[l] \sim [ɭ]$, $[\theta] \sim [f]$)¹. Three contrasts contained one member that is present in the Japanese inventory and one member that is absent ($[s] \sim [\theta]$, $[s] \sim [f]$, $[b] \sim [v]$). One contrast was chosen that is effectively a native contrast for Japanese learners as both members of the contrast are also present in the Japanese inventory ($[p] \sim [b]$).

The study was a pretest-posttest design in which two experimental groups of subjects received training once a week for five weeks beginning one week after the pretest and ending one week before the posttest. A control group was examined on the pretest and posttest with an interval of five weeks between them and received no training.

Subjects

A total of ninety-nine subjects participated in the study. Two experimental groups (group A, $n=27$; group B, $n=39$) received training in the careful articulation in the non-native segments under investigation. The two groups were from separate streams within the university, where they are registered in two different programs of the Faculty of Humanities at Hokkai Gakuen University. The subjects were, nevertheless, registered in separate sections of the same English courses at the time of testing and training. A third group of control

subjects (n=33) received no training. The members of the Control group were registered in the same two programs as the subjects in the experimental groups (n=15 and n=18, respectively)

Testing

Pretest and posttest consisted of the same set of stimulus pairs presented in an AX discrimination task. Subjects were instructed to indicate on a response sheet whether each pair of words they heard were instances of the same word or different words. Verbal instructions were given in English and accompanied by written instructions in both English and Japanese.

The stimuli were made up of real words spoken by one male voice with a standard North American accent. All of the lexical items were selected from materials used by the subjects in the vocabulary course they were registered in at the time of testing and training. There were twelve experimental pairs for each of the six contrasts used in testing ([s]~[θ], [s]~[f], [b]~[v], [l]~[ɹ], [θ]~[f], [p]~[b]). The “same” pairs contained non-identical instances of the same word. There was a pause of 1500 msec between members of each pair and 3000 msec between pairs. The stimuli were presented through Sony HS-90 headphones controlled through a Sony LLC-9000 control console.

Training

Five training sessions were administered once per week over a period of five weeks. Training began one week after the pretest and was completed one week before the posttest. Each session included training on all 5 non-native segments under investigation ([f], [v], [θ], [l], [ɹ]). The training methodology was neither multiple-talker nor single-talker but rather “no-talker”. Subjects received no perceptual

training at all. Consequently, training sessions did not include any auditory cues containing model pronunciations. Rather, subjects received explicit instruction in the precise articulation of the five non-native segments. The instruction was supported with silent, visual demonstrations of the articulation which were repeated as many times as the subject requested. Real words were silently articulated by the experimenter along with presentation of line drawing illustrations depicting the referents of the words. Subjects mimicked the experimenter's articulations silently, pronounced the words out loud, and received positive feedback when correct and continued correction with further demonstration when incorrect.

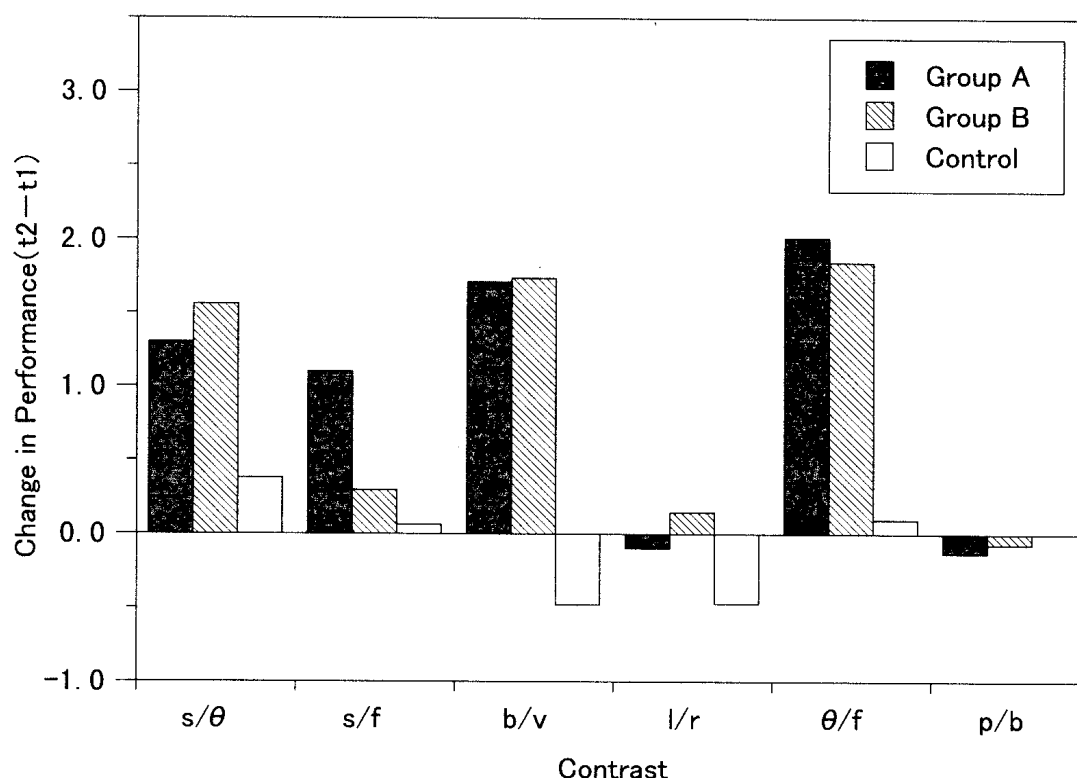
The motivation behind this somewhat radical approach was to avoid the development of stimulus-dependent representations that researchers using perceptual training have encountered (Lively, Pisoni, and Logan, 1992; Logan, Lively, and Pisoni, 1991). By providing subjects in this study with instruction and visual demonstration exclusively, any perceptual model they may have developed in the course of training could only come from their own articulations or the productions of other subjects in the same training session. The co-occurrence of their own articulations with the kinesthetic sense of the concomitant gestures was thought to reinforce the development of the new segmental category.

Results

Figure 1 illustrates the mean change in number correct from pretest (t1) to posttest (t2) on each of the contrasts under investigation as measured by subtracting pretest scores from posttest scores for the three groups of subjects.

The [p]~[b] and [l]~[ɭ] contrasts showed only negligible change

Figure 1. Mean change in number correct from pretest to posttest



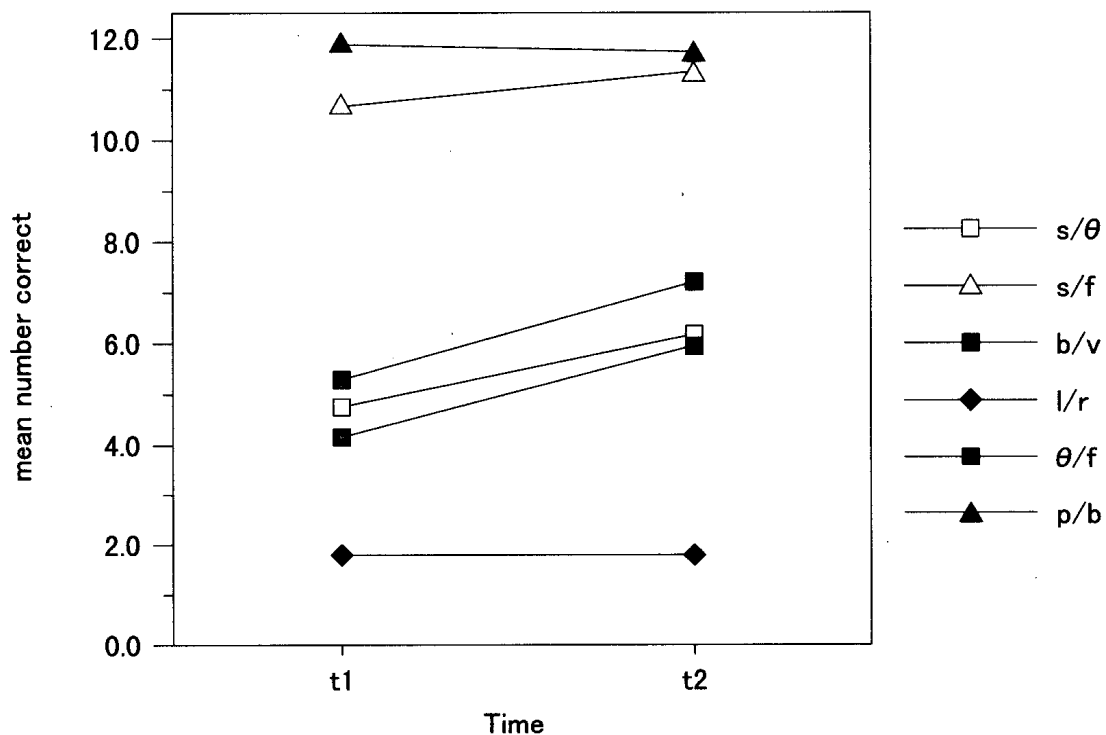
from pretest to posttest for subjects in all three groups. Under post-hoc analysis, *Scheffe F*-test, a factorial ANOVA reveals no statistically significant difference between any of the groups, $F(2,96)=.744$; $p=.478$ and $F(2,96)=1.517$; $p=.2247$, respectively. Although a difference between groups on the $[s] \sim [f]$ contrast did reach significance, $F(2,96)=3.209$; $p=.0448$, post-hoc analysis revealed that difference only to be significant between group A and the Control group on the Fisher PLSD, and no significant difference was measured according to the *Scheffe F*-test. Although the difference between groups did not reach significance for the $[s] \sim [\theta]$ contrast $F(2,96)=2.539$; $p=.0842$, there does appear to be a trend toward significance distinguishing the two experimental groups from the Control group. Significant differences were revealed in the changes from pretest to posttest on the $[b] \sim [v]$ and $[\theta] \sim [f]$ contrasts between group A and the Control group as well

as between group B and the Control group, $F(2,96)=10.61$; $p=.0001$ and $F(2,96)=9.994$; $p=.0001$, respectively. No difference was found between the two experimental groups.

As no significant difference was found between the two experimental groups (group A and group B), their scores were combined and submitted to a repeated measures ANOVA to determine whether there were significant differences between pretest and posttest scores on any of the contrasts under investigation. Figure 2 illustrates the mean number correct on the pretest and posttest for each contrast among the two experimental groups combined.

Post-hoc analysis, *Scheffe F*-test, reveals significant differences between pretest and posttest scores for the $[b] \sim [v]$ and $[\theta] \sim [f]$ contrasts only, $F(65,11)=1.037$; $p=.0001$. Although improvement from

Figure 2. Mean number correct for experimental groups on pretest (t1) and posttest (t2)



pretest to posttest on the [s]~[θ] contrast resembles that of both the [b]~[v] and [θ]~[f] contrasts, the difference did not reach significance. No difference was found between pretest and posttest scores for either the [p]~[b], [s]~[f], or [l]~[ɭ] contrasts.

Discussion

Although they received no perceptual training, subjects who received training in the pronunciation of non-native segments demonstrated significant improvement in their ability to discriminate those segments from other segmental categories upon perceptual testing. However, the positive effects of this training was not uniform across all contrasts studied. Table 2 indicates the six segmental contrasts investigated and how improvement following training related to the status of the segments in the L1 Japanese inventory.

Presence or absence in the L1 segmental inventory is not an adequate factor for determining the positive influence of pronunciation training on perceptual capacities. The two contrasts containing segments that are absent from the Japanese inventory (i.e., [l]~[ɭ], [θ]~[f]) differed with respect to the effects of training. The [θ]~[f]

Table 2. Relationship between improvement following training and the L1 Japanese inventory

	no improvement		improvement
	already acquired	not acquired	
both segments present in L1	[p] [b]		
one segment present in L1	[s] [f]		[b] [v] [s] [θ]
neither segment present in L1		[l] [ɭ]	[θ] [f]

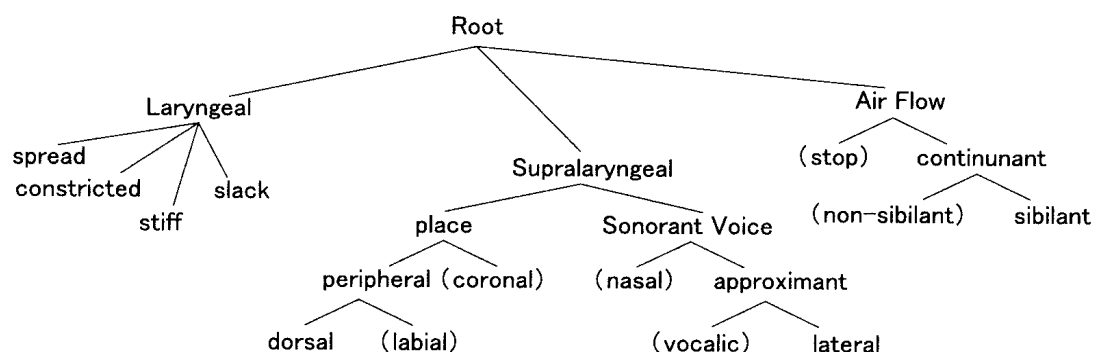
contrast showed significant improvement whereas the [l]~[ɭ] contrast did not. Moreover, the scores at pretest (as well as posttest) were considerably different. The [θ]~[f] contrast was correctly identified an average of 5.28 times out of 12. The segments [l] and [ɭ] were recognized as distinct segments an average of only 1.79 out of 12 trials.

The presence of one of the two segments in a contrast in the L1 inventory was similarly not sufficient to determine the effects of pronunciation training. Perception of the [s]~[ʃ] contrast showed no effect of pronunciation training, largely due to the fact that performance had already reached near ceiling levels at pretest (an average of 10.7 out of 12) leaving no room for improvement following training. Perception of the [b]~[v] and [s]~[θ] contrasts did improve following pronunciation training. Perception of the [p]~[b] contrast, which is a native contrast present in the L1 system, was trivially unaffected by training, as performance at pretest was already at ceiling (an average of 11.89 out of 12).

The relevant factor for determining whether second language learners can acquire novel segmental categories, I argue, is the system of segmental representations that they have developed in the course of acquiring their first language. Phonological theory provides a formal means of characterizing this knowledge within the framework of Feature Geometry. This approach to segmental representations recognizes sub-segmental features to be organized into a hierarchical structure that reflects relationships of dependency and constituency among the features. Figure 3 illustrates a geometry that is based on the one proposed by Rice (1992, 1995), though details distinguishing C-Place (for consonants) and V-Place (for vowels) features have been left out for reasons of perspicuity.

One important factor to note with this kind of feature geometry is

Figure 3. Feature Geometry



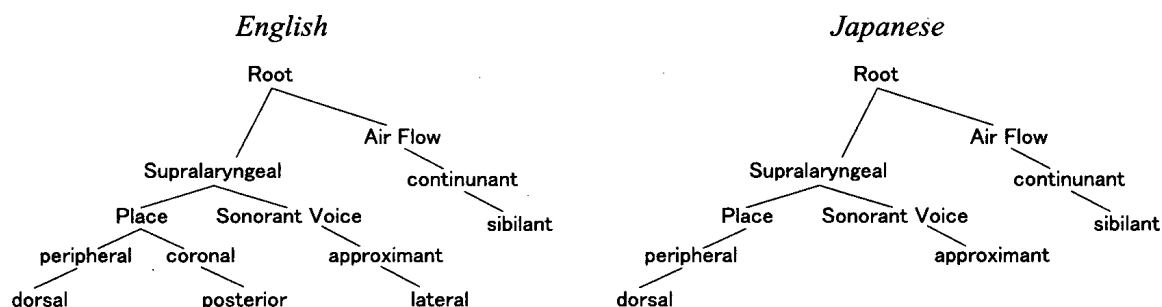
the status of features depicted within parentheses (e.g., (nasal), (labial)). These features are not present in underlying representations of individual segments but are, rather, interpreted by the absence of their sister feature. In the case of (labial), for example, the presence of a [peripheral] feature in a segmental representation with the absence of a subordinate [dorsal] feature is interpreted by the grammar to be a labial segment. Likewise, a segment containing a specification for [Sonorant Voice] but no subordinate [approximant] feature would be interpreted as a nasal segment. Substantial theoretical mileage is gained through this use of monovalency which is outside the scope of the present argument. What is relevant for us is the fact that though no language employs the full range of features represented in figure 3, all languages respect the specific dependency and constituency relations that are encoded in this geometry. Similarly, within a particular language, the individual segments that make up the segmental inventory do not exploit the full feature geometry. Only those features that are required to distinguish a segment from all other segments in the inventory are included in its segmental representation.

According to Brown and Matthews (forthcoming), the development of segmental structure in the course of first language acquisition proceeds by way of building structure in response to contrasts detected

in the input (see also Rice and Avery, 1995). A child begins with that amount of structure that is common to all languages and adds to that structure in a manner consistent with the dependency and constituency relations that are represented in figure 3. One child might elaborate structure under the Place node before structure under the Sonorant Voice node, for example, while another child might do the reverse. This building process continues until enough structure has been added to distinguish all contrasts in the system being acquired, at which point the child has arrived at the adult feature geometry.

The adult feature geometries for English and Japanese appear in figure 4 (the laryngeal features have been left off for ease of illustration). The structure that makes up the Japanese geometry is a subset of the English geometry. That is, the English system contains all of the phonological contrasts that the Japanese system has, plus additional contrasts that are not part of the Japanese system. Following Brown and Matthews (forthcoming), it is assumed that a child acquiring English could pass through a stage at which the emerging feature geometry would resemble the adult Japanese system. A child acquiring Japanese, however, would never pass through a stage at which the feature geometry resembled the adult English system since the Japanese input would never include contrasts to trigger building of the

Figure 4. Feature geometries of English and Japanese



additional structure.

Brown (to appear) has suggested that the feature geometry imposes on the perceptual system the specific boundaries within which phonemic categories are perceived. Therefore, once established, the segmental representation that defines a phonemic category will be activated by a range of phonetic variants within the category yet not be activated by instances of other categories in the system. Such organization is of considerable value in the processing of incoming speech since a certain degree of noise is permitted within categories without compromising the recoverability of the appropriate underlying form.

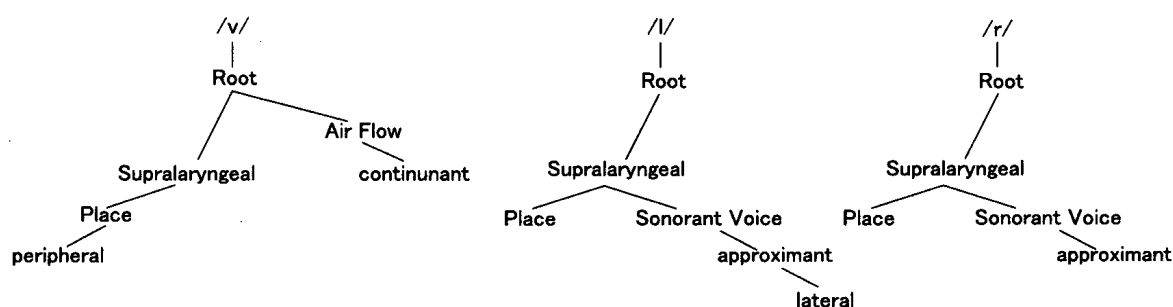
For the second language learner, acquiring a new phonemic category means acquiring a new segmental representation. However, the range of possible segmental representations that he or she will be able to construct will be constrained by the feature geometry that he or she has acquired in the course of first language acquisition. Because that feature geometry imposes organization over the acoustic perceptual system, the learner will only be sensitive to non-native contrasts that are distinguished along dimensions corresponding to features in the geometry. Contrasts within a native category — that is, non-native contrasts that would require new structure to be added to the feature geometry — are perceived as instances of the same category in much the same way as different phonetic variants of the same category in the native language are perceived. Without the necessary detection (or even perception) of contrastive use in the input to trigger the construction of novel segmental representations, the second language learner is unable to acquire the new category.

If, however, a non-native phonemic category is distinguished from other categories along a dimension that is already manipulated in the learner's native language, then he or she will be capable of perceiving

the contrastive use in the input and, therefore, ultimately be capable of acquiring the new category. Although Japanese lacks a voiced bilabial fricative /v/, for example, the Japanese feature geometry does contain the necessary features to build a representation for one. In contrast, the English /l/ and /r/ categories are distinguished by the presence or absence of the feature [lateral] under [approximant] in the feature geometry (cf. Brown, 1995). The Japanese system lacks this feature. Consequently, all approximants will be perceived by Japanese speakers simply as that, as approximants. The crucial acoustic differences between approximants will be treated as intracategory variation that can only be detected when presented in very controlled conditions (e.g., AX discrimination presentation with a very short intra-stimulus interval, on the order of 250 msec; see Werker and Logan, 1985). The segmental representations of English /v/, /l/, and /r/ appear in figure 5.

When a speaker of Japanese hears an English [v], he or she may perceive it as distinct from other categories; however, since there is no segment that corresponds identically to that particular constellation of features in his or her inventory, he or she may also experience interference from one or more of the native categories defined by native segmental representations. Through the course of second language acquisition, the learner will be able to construct a new segmental

Figure 5. Segmental representations of English /v/, /l/, and /r/



representation based on his or her detection of the contrastive use of that constellation of features in the input.

When a Japanese speaker hears an English [ɹ], he or she will recognize it as a non-nasal sonorant with no specification for place, corresponding roughly to the native category with that geometric representation, /ɾ/. When he or she hears an English [l], he or she will similarly recognize it as a non-nasal sonorant with no specification for place, corresponding to the same category. The presence of the feature [approximant] in the Japanese feature geometry with no dependents imposes on the acoustic perceptual system of a native speaker of Japanese the categorical perception of approximants with no category distinction among approximants.

A final bit of evidence in support of Brown's (to appear) proposal that the native feature geometry superimposes organization over the acoustic perceptual system comes from the effects of brain damage on the perception of non-native contrasts. Kohno and Tsushima (1997) have reported a study in which they compare the perceptual capacities of Japanese infants, Japanese adults and Japanese Broca's aphasics. Their findings replicated those of previous research conducted by these researchers which documents the capacity of Japanese infants to perceive the English [l]~[ɹ] contrast and the inability to do so among Japanese adults (Tsushima, Takizawa, Sasaki, Shiraki, Nishi, Kohno, Menyuk, and Best, 1994). What is of particular interest, however, was the performance of the group of aphasics. Some patients were able to perceive the English [l]~[ɹ] contrast significantly better than the unimpaired adults. Neurological damage to the phonological system in aphasia effectively impairs the operation of the feature geometry, thereby freeing the acoustic perceptual system from the feature geometry's superordinate organization over incoming linguistic stimuli and

permitting the individual to perceive the acoustic differences between those segments. Thus, without the intervening effects of the feature geometry, these aphasics have direct access to their acoustic perceptual system which has never lost the capacity to discriminate the [l]~[ɭ] contrast. This capacity has only been masked by the phonological system that developed in the course of first language acquisition and that has mediated perception throughout adulthood. This analysis echoes that of Miyawaki et al. (1975) who concluded that native language experience shapes the way listeners perceive the speech signal but does not change the underlying auditory capabilities of the listener.

Conclusion

The study presented here demonstrates that explicit instruction in the pronunciation of non-native segments can contribute to the development of novel segmental categories which can then be used to discriminate members of the novel category from members of other categories perceptually. However, not all non-native contrasts are created equal. Although pronunciation training can contribute to the development of new segmental representations, the phonological system of the native language constrains the process of that development. Specifically, the step-by-step acquisition of hierarchically organized segmental features in the course of first language acquisition results in the development of a mature feature geometry which subsequently imposes on the perceptual system the specific boundaries within which phonemic categories are perceived throughout life. When faced with non-native segmental categories, second language learners can only perceive them through their native language feature geometry. When the target segmental representation (i.e., the new segmental category) requires features that are not part of the L1 system, learners will be unable to perceive the

contrast in the input and therefore unable to acquire the new segmental representation. When the target segmental representation, though not in the native segmental inventory, can nevertheless be built from the set of features present in the L1 system, learners will be sensitive to the contrast between the new category and other categories. This sensitivity then triggers the construction of a novel segmental representation to underlie a robust new phonemic category.

Does this mean that Japanese learners of English will never be able to learn how to distinguish English [l] from [ɭ]? No. What it means is that Japanese learners of English will never be able to construct a segmental representation in the same way that native speakers of English do. It also means that a native speaker of Japanese who has not been exposed to spoken English from an early age will forever find it difficult to hear the difference between English [l] and [ɭ]. However, this need not be a particularly problematic issue since real-world communicative situations rarely arise in which a message is crucially determined by sensitivity to a single phonemic contrast. Top-down processing surely prevents a Japanese learner of English from misunderstanding a question like "Do you like to eat fried rice?" just as it prevents a native speaker from the same misunderstanding under noisy conditions. Moreover, learning to pronounce the difference between [l] and [ɭ] is well within the grasp of native speakers of Japanese. What is beyond their grasp is the capacity to reliably link those pronunciations to perceptual categories. This, also, need not be particularly problematic since orthographic information can be used very reliably to distinguish lexical items. As long as one can remember whether the word one wants to say is written with an "l" or an "r", one can learn to accurately articulate the difference between "playing in the lane" and "praying in the rain."

Thus, attaining proficiency in a second language need not correspond to developing a new linguistic system that parallels those acquired by native speakers of that language. Functionally effective means may be developed to communicate in a second language that may be qualitatively different from the means employed by native speakers. Indeed, this is the only recourse available when properties of the first language actually preclude a learner from acquiring properties of a second language.

Notes

*A version of this paper was presented at the New Sounds 97: Third International Symposium on the Acquisition of Second-Language Speech held at the University of Klagenfurt, Austria in September, 1997. Comments received there are gratefully acknowledged, though all remaining errors are, of course, my own.

¹It should be noted that Japanese does contain a segment with a voiceless bilabial fricative [ɸ] allophone that surfaces when followed by a high back vowel. Japanese does not have, however, a separate category for a voiceless fricative with labial articulation which might correspond to the English labiodental /f/.

References

- Brown, C. (to appear) The role of the L1 grammar in the L2 acquisition of segmental structure. *Second Language Research*.
- Brown, C. (1995) The feature geometry of lateral approximants and lateral fricatives. In H. van der Hulst and J. van de Weijer (eds.), *Leiden at Last*, pp.41-88. Leiden, Holland: University of Leiden Press.
- Brown, C. and Matthews, J. (forthcoming) The elaboration of segmental structure in first language acquisition. In S.J. Hannahs & M. Young-Scholten (eds.), *Focus on phonological acquisition*. Amsterdam: John Benjamins.

- Kohno, M. and Tsushima, T. (1997) Language-specific developmental changes in speech perception abilities and their neuropsychological reason — A hypothesis. Paper presented at the 5th Annual Congress of the International Society of Applied Psycholinguistics. Porto, Portugal: Universidade do Porto. June 25-27.
- Lively, S.E., Pisoni, D.B., and Logan, J.S. (1992) Some effects of training Japanese listeners to identify English /r/ and /l/. In Y. Tohkura, E. Vatikiotis-Bateson, and Y. Sagisaka (eds.), *Speech perception, production and linguistic structure*, pp.175-196. Tokyo: Ohmsha and Amsterdam: IOS Press.
- Logan, J.S., Lively, S.E., and Pisoni, D.B. (1991) Training Japanese listeners to identify /r/ and /l/. *Journal of the Acoustical Society of America* 89(2): 874-886.
- MacKain, K., Best, C., and Strange, W. (1981) Categorical perception of English /r/ and /l/ by Japanese bilinguals. *Applied Psycholinguistics* 2: 369-390.
- Mann, V.A. (1985) Distinguishing universal and language-dependent levels of speech perception: Evidence for Japanese listeners' perception of "l" and "r". *Cognition* 24: 169-196.
- Miyawaki, K., Strange, W., Verbrugge, R., Liberman, A., Jenkins, J., and Fujimura, O. (1975) An effect of linguistic experience: The discrimination of /r/ and /l/ by native speakers of Japanese and English. *Perception and Psychophysics* 18: 331-340.
- Mochizuki, M. (1981) The identification of /r/ and /l/ in natural and synthesized speech. *Journal of Phonetics* 9: 283-303.
- Rice, K. (1992) On deriving sonority: A structural account of sonority relationships. *Phonology* 9: 61-99.
- Rice, K. (1995) What is a lateral? Paper presented at the annual meeting of the Canadian Linguistics Association.
- Rice, K. & P. Avery. 1995. Variability in a deterministic model of language acquisition: A theory of segmental acquisition. In John Archibald, ed., *Phonological acquisition and phonological theory*, 23-42. Hillsdale, NJ: Lawrence Erlbaum Associates.
- Sheldon, A. and Strange, W. (1982) The acquisition of /r/ and /l/ by

Japanese learners of English: Evidence that speech production can precede speech perception. *Applied Psycholinguistics* 3: 243-261.

Tsushima, T., Takizawa, O., Sasaki, M., Shiraki, S., Nishi, K., Kohno, M., Menyuk, P., and Best, C. (1994) Discrimination of English /r-l/ and /w-y/ by Japanese infants at 6-12 months: Language specific developmental changes in speech perception abilities. Unpublished ms., Boston University, Boston, MA.

Werker, J.F. and Logan, J.S. (1985) Cross-language evidence for three factors in speech perception. *Perception & Psychophysics* 37: 35-44.