

Phonological features: Basic model

I. Why use phonological features?

We have found that some phonetic properties of sounds are relevant for their phonological behavior. For example, groups of sounds with shared properties often behave as a group in the phonological grammar, as seen in the Arabic problem. (A group of sounds that behave as a group like this is called a *phonologically active class*.) Another example of the importance of the phonetic properties of a sound is found in the common process known as *assimilation*, in which one sound takes on some or all of the properties of a neighboring sound.

We propose that in the mental grammar, *a segment is actually represented as a collection of properties*. That is, *it is the properties* (NOT segment-sized units) that the mental grammar is able to identify, store, and manipulate. This explains why sounds pattern in classes: the mental grammar deals in sound properties, and properties identify classes of sounds.

So, we have determined that phonetic properties of sounds are useful and important for understanding phonological patterns, and we propose that properties are the units with which the mental grammar operates. But not all aspects of the phonetics of a sound are cognitively — which is to say *phonologically* — relevant. The phonology does not seem to be sensitive to properties like the total amount of energy expended in producing a certain sound, or the distance in millimeters that the tongue is displaced, for example. So, *our phonological model needs to determine which phonetic properties are relevant for phonology and which are not*.

Model construction: We start by proposing that phonologically relevant sound properties are included in our model of the phonological grammar as *phonological features*. Many of these features are *binary*, meaning that they have a [+] value (for sounds that have a certain property) and a [–] value (for sounds that do not). Other features are *privative*, also called *monovalent* or *unary*, meaning that they are simply present or absent without a [±] value. In our phonological model, *a segment is represented as a bundle of phonological features*.

What is the justification for claiming that some property of a sound (phonetic, or otherwise) should be given the status of a phonological feature in our model?

- The most important criterion is that the property in question be needed to *account for some aspect of phonological behavior* in our data: for example, it is needed to define a phonologically active class, or it is explicitly changed in a phonological process.
- Another justification for introducing a new phonological feature is the need to *distinguish two sounds* that are treated differently in some language but are not already distinguished by phonological features in any other way.
- Ideally, all phonological features will also have a well-motivated *phonetic basis*, which is hypothesized to help in language acquisition. But in the model we are proposing, the first two criteria are the most essential.

The general proposal that features are part of the mental grammar is accepted by almost all phonologists. However, we also need to determine *what the features are*, and *how they are defined*. Here, there is a lot more controversy among phonologists, because it turns out to be very difficult to develop a set of features that works exactly as desired for all of the languages whose phonologies have been studied; it may even be that not all features are universal. We will consider some of these debates as we examine data from various languages. For this class, we will start out by *proposing the following feature set*, and we will use this set of features unless and until we *make an explicit argument* that it should be changed.

II. Our model of phonological features

The features listed below are each given:

- a phonetically based definition (as a memory aid — *not a formal part of our model*)
- **a description in terms of which natural classes of sounds they distinguish (this is the actual content of our feature model!)**

In learning these features, focus on how they designate, or distinguish between, **sound classes**.

A. Features specified for both consonants and vowels

[±consonantal]

Phonetic basis: [+cons] segments have at least as much constriction in the vocal tract as a liquid. [–cons] segments do not.

Model construction: Distinguishes vowels and glides from non-glide consonants:

[+cons] — Stops, fricatives, affricates, nasals, liquids

[–cons] — Glides, vowels

- Glottal segments such as [h] and glottal stop pattern phonologically as [–cons] in some languages, probably because their only constriction is right at the glottis.

[±syllabic]

This feature does not have a phonetic basis! Descriptively: [+syll] segments form the *nucleus*, also called *peak*, of a syllable. [–syll] segments do not.

Model construction: Distinguishes vowels from glides; also distinguishes syllabic consonants from other consonants in languages where this distinction is relevant.

[+syll] — Vowels and syllabic consonants

[–syll] — Glides and non-syllabic consonants

- Note that this feature definition is not based on *phonetics*, but purely on *phonology*.

[±sonorant]

Phonetic basis: [+son] segments have frictionless airflow in *either* the oral or the nasal tract (so nasal stops are [+son]). [-son] segments have airflow that is significantly obstructed in the vocal tract *overall*.

Model construction: Distinguishes sonorants from obstruents:

[+son] — Sonorants (nasals, liquids, glides, vowels)

[-son] — Obstruents (stops, fricatives, affricates)

[±voice]

Phonetic basis: [+voi] segments are produced with vibrating vocal folds. [-voi] segments are not.

Model construction: Distinguishes voiced segments from voiceless segments:

[+voi] — Voiced segments

[-voi] — Voiceless segments

[±continuant]

Phonetic basis: [+cont] segments are produced with moving air in the *oral* tract. [-cont] segments are not. (Note that the status of the *nasal* tract is not relevant for this feature.)

Model construction: Distinguishes stops — oral *and* nasal — from all other segments:

[+cont] — All segments other than oral and nasal stops

[-cont] — Oral and nasal stops

- What about affricates? They are potentially “both” [+cont] and [-cont], since they are phonetically a stop+fricative. The best way to address this question is to see how affricates pattern in the language that you are working on. Do they form natural classes with stops, or with fricatives? The English postalveolar affricates are often proposed to be [-cont].
- Liquids can also vary from language to language with respect to [±cont]. Laterals have moving air at the side(s) of the oral tract, so they are often [+cont]. But in some languages, laterals pattern phonologically as [-cont]; this is also phonetically plausible, because they do have a complete constriction in the center of the oral tract. Similarly, taps and trills have a complete constriction, but it is very brief (for a tap) or interrupted (for a trill), so languages vary in how they assign [±cont] values to these segment types. (On the other hand, [ɹ] as found in many English varieties is articulated much like a glide, so it is reliably [+cont].)

[±nasal]

Phonetic basis: [+nas] segments have a lowered velum, which allows nasal airflow. [-nas] segments do not.

Model construction: Distinguishes nasal segments from oral segments:

[+nas] — Nasal stops, nasalized vowels, other nasalized segments

[-nas] — All oral segments

[±strident]

Phonetic basis: [+strid] segments are produced with high-frequency fricative noise. [–strid] segments are not (they may be fricatives/affricates with lower-frequency noise, or they may be non-fricatives).

Model construction: For coronal segments, distinguishes sibilants from nonsibilants. Some phonologists also use this feature to distinguish labiodentals from bilabials, but this is more controversial.

[+strid] — Alveolar and postalveolar fricatives and affricates (and for some phonologists, labiodental fricatives and affricates)

[–strid] — All other segments

[±lateral]

Phonetic basis: [+lat] segments are produced with lateral (side) airflow around a central constriction. [–lat] segments do not have exclusively lateral airflow (they may have central airflow, or none).

Model construction: Distinguishes lateral segments from central segments:

[+lat] — Lateral segments, including [l], the palatal lateral liquid in Spanish, the coronal lateral fricative in Welsh, etc.

[–lat] — Central segments (all segments that are not lateral)

B. PLACE FEATURES: These features are specified for consonants.

Note that the so-called major place features are *privative*, written in small-caps by convention. (What difference does privative vs. binary entail in the predictions our model makes for the phonological behavior of a feature?)

If a consonant has two simultaneous distinct places of articulation, such as a labiovelar consonant, it has two place features at the same time.

[LABIAL]

Phonetic basis: [LAB] segments have the lower lip as their active articulator.

Model construction: Identifies labial segments:

[LAB] — Bilabial and labiodental consonants

[CORONAL]

Phonetic basis: [COR] segments have the tip or blade of the tongue as their active articulator.

Model construction: Identifies coronal segments:

[COR] — Dental, alveolar, postalveolar, and retroflex consonants; palatals are [COR, DORS]

[±anterior]

Phonetic basis: [+ant] segments are produced in the forward half of the coronal region; [–ant] segments are produced in the posterior half.

Model construction: Distinguishes between anterior and posterior coronal segments:

[+ant] – Dental and alveolar consonants

[–ant] – Postalveolar and retroflex consonants

- This feature is special because only segments that are [CORONAL] have any value for [±ant] at all. [COR] segments may be [+ant] or [–ant], but other segments are *neither*.
- The American English rhotic (IPA symbol: [ɹ]) is arguably [–ant]. However, non-retroflex flaps and trills are usually dental or alveolar, and therefore [+ant].

[DORSAL]

Phonetic basis: [DOR(s)] segments have the back of the tongue body as their active articulator.

Model construction: Identifies dorsal segments:

[DOR(s)] – Velar and uvular consonants; palatals are [COR, DORS]

- We have not yet considered how to **distinguish** velar from uvular consonants.

[GLOTTAL]

Phonetic basis: [GLOT] segments have the glottis (vocal folds) as their active articulator.

Model construction: Identifies glottal segments:

[GLOT] – Glottal consonants

C. VOWEL FEATURES: These features are specified for vowels.

[±high]

Phonetic basis: [+hi] segments have the tongue body higher than neutral (mid) position; [–hi] segments do not.

Model construction: Distinguishes high vowels from non-high vowels:

[+hi] – High vowels

[–hi] – Mid and low vowels

[±low]

Phonetic basis: [+lo] segments have the tongue body lower than neutral (mid) position; [–lo] segments do not.

Model construction: Distinguishes low vowels from non-low vowels:

[+lo] – Low vowels

[–lo] – High and mid vowels

- Mid vowels are [–hi, –lo]
- No segment can be simultaneously [+hi] and [+lo]

[±back]

Phonetic basis: [−bk] segments have the tongue body farther forward than the neutral (central) position; [+bk] segments do not.

Model construction: Distinguishes non-front vowels from front vowels:

[+bk] — Back and central vowels

[−bk] — Front vowels

- Our feature system does not easily distinguish between central and back vowels that are the same with respect to height and rounding (except potentially through a difference in [±ATR]; see below). This is not an accident — many phonologists have argued that languages never, or almost never, distinguish central and back vowels of the same height without some additional difference, such as length, rounding, or ATR. Therefore, it is reasonable to give them the same *phonological* featural representation, even though they are *phonetically* distinct.

[±round]

Phonetic basis: [+rd] segments have lip rounding; [−rd] segments do not.

Model construction: Distinguishes round vowels from unrounded vowels:

[+rd] — Round vowels

[−rd] — Unrounded vowels

[±ATR] (advanced tongue root)

Phonetic basis: [+ATR] segments have the root of the tongue in an advanced (forward) position; [−ATR] segments do not.

Model construction: Distinguishes tense vowels from lax vowels:

[+ATR] — Tense vowels, including [i e y u o]

[−ATR] — Lax vowels

- There is some debate over the nature of this feature in the languages of the world. We will assume that the [±ATR] feature involved in vowel harmony in languages like Akan is phonologically the same feature that distinguishes “tense” and “lax” vowels in English, but it may have slightly different phonetics in different languages.
- In languages with a traditional “tense/lax” distinction, [+ATR] corresponds to “tense,” that is, to a more extreme/less mid¢ral position in the vowel space.
- *Phonetically* similar vowels may be *phonologically* classified with different ATR values in different languages, especially for central vowels and low vowels. Looking to see what other vowels a particular vowel forms phonologically active classes with may help determine its value for [±ATR] in a given language.
- In many languages with small vowel systems, [±ATR] is not active; that is, it is never phonologically relevant (and can therefore be ignored in a phonological analysis). A useful rule of thumb: *Don’t add [±ATR] to your analysis unless patterns in the data show that you need to.*

III. Using features to describe segment and segment classes

When we use our phonological model to describe a segment or a segment class, we write all of the features that belong to the same segment inside one set of brackets. Technically the model uses a vertical format, as on the left below, but we can use the horizontal format with commas, as on the right, because it is easier to type.

- For example, the description “labial oral segments” would be represented as:

$$\begin{bmatrix} \text{LAB} \\ \text{-nas} \end{bmatrix} \quad \text{or} \quad [\text{LAB}, \text{-nas}]$$

- But those feature representations mean something different from this:

$$[\text{LAB}] [\text{-nas}]$$

which represents a sequence of two segments where the first one is labial, and the second one is oral.

What is the *best way to describe a segment class* in our model? Here are some important strategies:

- Usually we *use as few features as possible*, while still accurately distinguishing the class we are interested in from the segments that must be excluded. Having our analysis refer only to the necessary features helps us determine which features really matter for modeling (understanding) a given phenomenon.
- Sometimes we choose which features to specify based on *what best helps us describe, predict, or explain the phenomena we are analyzing*. For example, should we describe the segment class [u o] as [+bk, +rd] or [+bk, -lo]? This might depend on what other segments are present in the language, or on the nature of the phenomenon that this segment class is participating in.