

Originally delivered in abridged form at the 2009 annual conference of the Society for Ethnomusicology, November 22nd, 2009; Mexico City, Mexico
Sami Abu Shumays, Independent Scholar

The Fuzzy Boundaries of Intonation in *Maqam*: Cognitive and Linguistic Approaches

“Nearly all this variation in context and sound comes from different ways of dividing up the octave and, in virtually every case we know of, dividing it up into no more than twelve tones. Although it has been claimed that Indian and Arab-Persian music uses “microtuning”—scales with intervals much smaller than a semitone—close analysis reveals that their scales also rely on twelve or fewer tones and the others are simply expressive variations, glissandos (continuous slides from one tone to another), and momentary passing tones, similar to the American blues tradition of sliding into a note for emotional purposes.” (Daniel Levitin, *This is Your Brain on Music*, New York: Plume, 2006. p. 39)

So claims Daniel Levitin in his otherwise excellent and insightful book, *This is Your Brain on Music*, which synthesizes and explains a great deal of the recent work on Music Cognition. Anyone with an amateur-level interest in World and Traditional musics, let alone Ethnomusicologists, knows the above statements are patently false—both his statement about Arabic, Persian, and Indian music, as well as his statement about American Blues music. So why does such a claim make it into such an important book? The answer becomes clear on delving deeper into the book, or into other work that has been done on the cognition of music: almost all of the research done into how the brain processes music has dealt exclusively with Western Classical music, and cognitive scientists who are interested in music have little or no familiarity with the music of the rest of the world.

I believe Arabic music has many features that can contribute to a compelling discussion on the cognitive science of music, not least because of its rich microtonal intonation system. The central argument of this paper is that music shares many features with spoken language as an information processing, storage, and communication system in human cognition—I will examine some of these parallels from the perspective of the *maqam* tradition, the modal melodic system of improvisation and composition that forms the basis of music in the Arab world, and which extends in variant forms from North Africa through Western and Central Asia.

Intonation will be the focus of my argument; and while I don’t need to convince ethnomusicologists of the existence of notes in Arabic music (and the music of other *maqam* traditions) beyond the 12 tones of an even-tempered Western scale, I am aware that the lack of an adequate explanation for the tunings of those notes in *maqam* scales is an important factor allowing music theorists to ignore their existence, and scientists who consult those theorists to be unaware of the tunings altogether—much as American music theorists, lacking an adequate explanation for the microtones used in Blues scales, mis-heard and miscategorized those notes as “expressive variations, glissandos (continuous slides from one tone to another), and momentary passing tones... sliding into a note for emotional purposes,” and passed such explanations on to Levitin and others.

Problems with Current Tuning Schemas

In advancing a better explanation for the microtonal tunings of *maqam* scales used in Arabic Music, it will be helpful to begin with the various explanations that have been posited to date. The first such is the hypothesis that Arabic scales are based on 24 equal divisions of the octave—an

explanation given credence by the fact that Arab musicians use the term “quarter-tone” when describing some of the notes of their scales. Yet this hypothesis was disproven at the 1932 Cairo conference on Arabic Music, when theorists tuned a *qanun* (a 72-stringed zither) to a 24-tone scale, and then asked Arab musicians to play on it. Those musicians universally found that scale to be out of tune, and retuned the *qanun* to their liking. Despite the tendency of some theorists to deny acoustical reality in favor of an explanation that better fits their theory, the theorists at the conference accepted this fact as a disproof of the 24-equal-divisions hypothesis.

Yet there is another very good reason such an intonation system is unlikely: Arabic music is an oral tradition, and equal divisions of the octave require a complex process to generate (and irrational numbers based around roots of the number 2), because of their contortion of the natural acoustical intervals—so there would be no reason to develop such a system without the kinds of modulation and harmony in Western music that led to the need for equal temperament, and no way to implement such a system without the kinds of instruments used to do so in the West (we can safely assume that the music theorists who retuned a *qanun* to 24 quarter tones in 1932 did not do so by ear, but rather by consulting a device that was able to generate frequencies with ratios of perfectly even 24th roots of the number 2).

This same argument applies to several of the other theories advanced to explain the failure of the 24-division schema: namely, theories about 17 or 22 equal divisions of the octave, as well as the popular “comma” theory of 53 or 54 equal divisions of the octave (9 equal divisions of a whole step; $9 \times 6 = 54$), used in Turkey to explain their scales; I have encountered musicians in Syria and Egypt using this explanation as well. *An equal-divisions schema doesn't make sense as the explanation for intonation in an oral tradition, because it doesn't account for how those intervals are generated in that tradition.* Beyond that, the use of Just 4ths and 5ths, which cannot be achieved through *any* equal division schema, in the tuning of musical instruments in the Middle-East, should be as conclusive a proof as any that such an approach is the wrong one to take.

So do intervals in Arab *maqam* scales correspond to some other acoustical reality? Certainly we can explain many of the notes of Arabic scales in terms of Just intonation: because of the presence of stringed instruments tuned in fourths and fifths, and the importance of fourths and fifths to the overall structure of Arabic scales (which are built around tetrachords and pentachords), many of the notes in Arabic scales follow some kind of Pythagorean tuning. Can we explain the microtones in terms of Just intonation as well? Some have tried: as early as the 8th century, Zalzal, an important musician and theorist from Baghdad, hypothesized that the “neutral third” (the third in between the major and minor third in size) had a frequency ratio of 27/22, which seems a bit complex until we realize that its inverse in the perfect fifth is 11/9 ($27/22 \times 11/9 = 3/2$), the ratio between the 11th and 9th partials of a tone. Such an explanation fits within the theoretical requirement, taken originally from the ancient Greeks (whom the Arabs studied in detail), that harmonious intervals are made up of small-integer ratios—and a trained musician can both hear and generate the 9th and 11th partials of a tone. So is this enough to account for the notes in Arabic music?

In my own practice as an Arab musician, I've identified at least 12 notes between my lowest e-flat and my highest e-natural (comparing the different *maqam* scales alongside each other such that the variable notes occur on some kind of E: low-e-flat in Bayati on G, e-flat in Nahawand on C, e-flat in Kurd on D, e-flat in Hijaz on D, e-half-flat in Saba on D, e-half-flat in Bayati on D, e-half-flat in Rast on C, e-half-flat in Sikah on e-half-flat, e-half-flat in Sikah Beladi on C, low-e-natural in Jiharkah on C, e-natural in Hijaz on C, e-natural in Ajam on C).¹ This is within one regional practice (my playing is Egyptian, despite having spent time in both Syria and Egypt); some of those same

¹ http://media.libsyn.com/media/shumays/PPerform_019_2007-02-02.mp3 for my demonstration of this.

notes are tuned higher in Syria than in Egypt, in particular the notes identified as e-half-flat. Scott Marcus gives a description of this phenomenon in Egypt: that many (though not all) musicians use slightly different intonations for their microtones depending on the *maqam* in which they occur; in his description he accounts for several of the microtones I use myself.

Now, that being the case, it should start to become clear that the other principal Western theoretical tool for understanding intonation, small integer ratios of frequencies, couldn't possibly apply. There are not enough such small integer ratios to account for these 12 E's, if we treat them as different kinds of thirds from C. Here is a list of all of the ratios of small integers lying between the Pythagorean minor third and the Pythagorean major third:

$32/27 = 1.185185\dots$ (*the Pythagorean minor third*)

$19/16 = 1.1875$

$6/5 = 1.20$ (*the Just minor third*)

$17/14 = 1.2142857\dots$

$11/9 = 1.2222\dots$

$27/22 = 1.2272727\dots$

$16/13 = 1.2307692\dots$

$5/4 = 1.25$ (*the Just major third*)

$81/64 = 1.264625$ (*the Pythagorean major third*)

I have treated frequency ratios using prime numbers above 20 as too difficult to perceive distinctly as such; the higher (non-prime) integers in the ratios of the Pythagorean major and minor thirds do not render those thirds unhearable or unplayable, since those more complex ratios like $32/27$ are perceived as ratios built from other common ratios: in this case, $8/9$ of $4/3$, or a Just whole step below a Just 4th. But some of these in the list above are arguably too difficult to hear (like $17/14$ in particular), to serve as a convincing basis for some of the Arabic intonations—why and how could we perceive a tonic note as the 14th partial of some other note, such that we could then hear the 17th partial above that same fundamental in relation to the 14th partial, as some kind of third? Such a ratio wouldn't be treated as harmonic by the ear without an *intentional* process generating that fundamental and training the ear to hear it; I was never taught such a process in order to achieve the correct intonation of microtones, and yet I play and sing them in tune.

But even if we were to accept some of these ratios as valid for the sake of argument, we still haven't accounted for the 12 Es I play, or for the additional Es within that range in the Syrian tradition. So while such ratios might account for *some* of the Es used in Arabic music (in particular the strong harmonic resonance of the Pythagorean and Just major and minor thirds), in no way can we account for all of them in this way. (This is another argument against the “comma” theory, which proposes 9 equal divisions of a whole step—since there are more than 12 distinct notes within a half-step in practice.)

How can we suitably define these intervals, then? Certainly we can measure the practice of musicians and give a cent value to the intervals they use. But would those cent values be an explanation for why those intervals are the way they are? The problem is the same as in the case of frequency ratios such as $17/14$, and the case of equal divisions of the octave: any claim about the intonation of a particular note in a musical scale *has to include information about why and how that intonation is used by practicing musicians*—which is certainly true of all of the intonation theories describing Western music². So does telling you that this E is 143 cents above D give you enough

² Just intervals are both audible and reproducible by the voice, and on string and wind instruments; the more complex forms of temperament developed for European keyboard music from the 15th century on were tunable by counting beats, and they served specific needs of practicing musicians.

information to be able to reproduce it? No, not without the aid of an electronic tone generator—yet practicing musicians can reproduce these very fine intonations perfectly, and audiences familiar with the music can recognize even slight deviations from the correct intonation.

In sum, neither Just intonation nor Equal-Temperament schemas can account for the microtonal intonations used in the Arabic *maqam* tradition. We can certainly understand why the failure of Western temperament schemas to account for Arabic music might lead Western theorists to deny the validity or correctness of such intonation, and to claim instead that Arabic music is out of tune, or that the intonations used are “mistakes,” or simply “expressive variations, glissandos, and momentary passing tones.” We can also understand why Arabs theorizing about their own music using the Western tools available to them might try to make *corrections* to their intonation based on those Western schemas, as happened throughout the 20th century—but this is not the place for a broader discussion of the effects of Western colonial and cultural hegemony on Arab music theory and practice, and on Arab perceptions of their own musical practice as somehow inferior and in need of correction. By and large, however, such theoretical discussions and attempts have not affected the intonation of practicing musicians in the Arab world—and the status quo is a seemingly unbridgeable gap between theory and practice that has left many Arabs to feel that their own music lacks the precision and correctness of Western music.

This is where knowledge of how a tradition is *learned* can begin to inform our analysis of its structure. Microtonal intonations presently used in Arabic music can be learned in only one way: by listening to them over and over again, and by imitating and repeating them back. But when students and musicians imitate, they don’t simply repeat one note in isolation. Rather, they repeat whole melodies. Only melodies give the particular feel and sound of intervals, because of the relationships with surrounding notes. I learned that as a student, myself: that to play the microtonal notes used in *maqam* scales in tune, I had to play melodies and melody fragments that I had already learned by ear, containing those notes. Without learning those melodies, there was no way to play those notes in tune consistently.

Melody Words and Musical Vocabulary

In exploring this phenomenon, the first and most fundamental parallel I’d like to make between music and language is that the primary structural unit of music is the *word*, rather than the individual note, in nearly the same sense that a word is the primary structural unit of language. In spoken language, words are made up of sequences of a limited number of possible phonemes, and in music words are made up of sequences of a limited number of possible notes. It has been widely acknowledged in Ethnomusicology that elements variously called “motifs,” “formulas,” “patterns,” “phrases,” etc. form the basis of oral composition and improvisation, at least since Albert Lord published *A Singer of Tales* in 1960 (Hill: 97). Cognitive science has shown that grouping, or “chunking,” of elements is a common feature of many areas of perception and cognition (Pinker 1997; Lerdahl and Jackendoff 1983: 13; Levitin 2006: 77). Such an understanding, though superior to a note-by-note analysis of music, is not rigorous enough for a deeper understanding of musical structure and cognition. Let me illustrate what I mean in terms of Arabic Music: If we take a long melody in *Maqam Rast*, such as:

Example 1



we find that we can break it up, or *parse* it, into several shorter phrases:

Example 2



Now, the reality for me is that I don't process that long phrase as a sequence of 27 notes, but as a sequence of 5 shorter phrases. But it is not simply that I group that melody into those phrases *in a particular moment*, based on the prevailing musical texture or emphasis—which is the sense in which Lerdahl and Jackendoff intend grouping in musical structure—but rather that I identify each one of those phrases as a distinct entity *that I have stored in memory*. I do not spontaneously generate those shorter phrases from a set of rules about how notes go together in the *maqam* system—I have learned them over years of imitating and repeating, in the same way that I learned the words of spoken language, and each one has its own identity. Words must each be individually learned and stored in memory, and can be recognized as an identifiable and remembered unit by other speakers of the language, just as we remember and identify faces, names, and many other common elements of daily cognition.

This is the more rigorous sense in which I use the term “word” to account for the unit of musical structure: a recognizable and reproducible combination of notes that is stored in my memory and in the memory of other musicians and listeners in my tradition. Levitin discusses some experimental evidence regarding varying musical perception of notes, depending on where they occur within a phrase—in those experiments, notes are perceived most strongly at phrase boundaries (158)—but he and the original experimenters don't come to the same conclusion I have because of their Western-music-theoretical blinders (they're not looking for or expecting to see melodic words); Levitin claims instead that “we expect certain pitches, rhythms, timbres, and so on to co-occur based on a *statistical analysis* our brain has performed of how often they have gone together in the past” (114-5, emphasis mine), with no explanation of why the brain would employ such a complicated statistical mechanism rather than the much simpler mechanism of storing combinations of pitches or rhythms in memory as words. The storage of musical words in memory is therefore an important area for exploration in the cognitive science of music, but as a practicing musician, I take it for granted—as do other practicing musicians of oral traditions, even if they do not express it in these terms (many do, however). As soon as we recognize words, we begin to recognize the existence of a finite but large vocabulary stored in memory, from which musical structure is generated. Once we've heard one of these melody-words enough times and in enough combinations with other melody-words, we begin to remember and recognize it whenever we hear it—in other words, we begin to parse it.

Let's take another melody-word in *maqam rast*:

Example 3



and see how it can occur in another long phrase (indicated with bracket):

Example 4



and in yet another phrase:

Example 5



If you look carefully at examples 4 and 5, you will see that I have used several of the melody-words from example 2 as well as the one from example 3.

Now we may characterize one of the basic structural similarities of Music and Language: they both are *discrete combinatorial* systems, a type of system in which a finite number of definable (discrete) elements combine in different ways to produce potentially infinite possible combinations (Pinker: 76). In any given language, there are an infinite number of possible sentences that can be constructed, out of a finite number of distinct words. It is not necessary to store all possible sentences in memory (which would be physically impossible, as it would require infinite storage space); rather, by storing the finite number of rules and principles of word combination, it is possible to generate infinite variety from finite means.

But language is not only a discrete combinatorial system in terms of words organizing into sentences, it is also a discrete combinatorial system in terms of phonemes organizing into words. From a small number of phonemes (between around forty and around one hundred and sixty for all human languages), it is possible to generate an infinite number of words. It is true that we find a finite number of words in the language, but new words are invented all the time and added to the language—sometimes they stick and sometimes they don't, but there is still a potential infinity of words, generated from those phonemes used in the language. This important characteristic of having discrete combinatorial systems operating on two levels—with the bulk of memory storage going to the element at the middle level of those two systems (words, which are in between phonemes and sentences)—is something music shares with language as well. Any given musical system has a small number of basic sounds—the tones and intervals of the scale, the smallest units of ornamentation (trills, turns, vibratos)—which can combine in potentially infinite ways, but which produce a large but finite vocabulary of melody-words; words which in turn can combine to produce an infinite number of longer melodies. In music, new melody-words may be created or invented much more frequently than in spoken language, but that doesn't diminish the fact that we still perceive them as distinct chunks, remember them when they reoccur (if we are listeners), and imitate them (if we are musicians). And at any given point in time, in the practice of a given musician, there are a finite (but extremely large) number of such words.

My point is that we cannot simply jump from the basic “phonemes” of the music—the notes, intervals, and ornaments—to the infinite potential melodies we find in music: some examples of basic phoneme combinations simply don't occur in the melodies of a given genre of music. If we start with only the notes of the Western diatonic scale as building blocks (which are used in several of the Arabic maqamat that don't include microtones), some possible combinations of them sound like Arabic music:

Example 6



some sound like 18th century Western music—like this example taken from a Mozart Piano Sonata:

Example 7



and some sound like late 19th century Western music, like this example taken from my favorite movement of a Brahms Symphony:

Example 8



But If I take *words* from 18th century music, I can combine them in an infinite number of ways, but I will always get new melodies that still sound like 18th century music, no matter how I combine them, and which will never sound like Arabic melodies or like Brahms (except where Brahms' vocabulary overlaps with Mozart's—and we should expect some overlap—or where he is intentionally evoking an antique style). Here's my newly composed example from snippets of familiar Mozart melodies I have in my melodic vocabulary, from years of studying Western music and playing the piano; you can see the recombination of some melody-words from the previous Mozart example:

Example 9



These examples have, in the case of Arabic music, used only the notes (the phonemes) that it shares with Western music. Of course, if we use the microtonal intervals and particular ornaments unique to Arabic music, we will generate something that sounds like Arabic music and not Western music, but I wanted to make the point that it is the finite vocabulary of *combinations* of notes into words, even more than simply the notes themselves, that gives the character of a particular kind of music.

General results in cognitive science, when weighed alongside the evidence from ethnomusicology, suggest strongly, I would argue, that there must be an *innate* mechanism for assembling music into word-like units. First, one thing that has become clear in cognitive science over the last few decades is that *there is no such thing as general intelligence*, from which humans develop their various cognitive abilities (Pinker 1997). Every cognitive ability has some kind of module in the brain allowing for it, from recognizing faces to motor control to the many different distinct modules involved in language processing, including *separate* modules for processing and parsing sound, for evaluating grammar, for connecting words to meaning, etc.

The existence of a discrete combinatorial system operating on multiple levels, and the ability to store and recall melody fragments from memory, suggests that there must be specific innate mechanisms in the brain for grouping musical sounds and storing them as discrete units. We don't, for example, create a discrete combinatorial system from animals—even if we see lions, zebras, and giraffes together in the wild, and then in the zoo, and then in the wild again, and then in another zoo, enough times to stick in our memory (the number of times, in other words, that a child hears a word before she starts to perceive it as a word rather than just simply a sequence of phonemes), we will not begin to think of lion-zebra-giraffe as a single inseparable unit and giraffe-zebra-gazelle as

another such memorizable unit. So discrete combinatorial systems do not simply arise spontaneously from combining any elements of cognition—rather they arise in the specific instances in which they are useful. Information theory has already demonstrated the vastly greater efficiency of coding information into words rather than just letters in written language (Pierce 1980); such results are generalizable to any discrete combinatorial system in signal transmission; we do not have room to discuss those results in detail here—suffice it to say that from the perspective of an economical use of resources in the brain, we should expect to find *any* signal system organized into words rather than simply base-level phoneme-like units. As I’ve hinted above, this is one of the areas where the almost exclusive focus of research in the cognitive science of music on Western music has prevented researchers from asking the right questions, because the orally-transmitted linguistic elements within Western music itself have been obscured as the result of centuries of bias toward learning from, and analysis of, notated music.

Returning to the issue of microtonal intonation in Arabic music, we find as practitioners that the intonations of particular intervals in Arabic music are inseparably embedded in melody-words in our memory, just as in spoken language our perception of particular phonemes is inseparable from the common words in which they occur, a fact demonstrated easily when we use whole words to illustrate and teach the pronunciation of phonemes: e.g. “ah” as in “father” vs. “a” as in “fatter.” In practice, there is no good way to teach phoneme pronunciation in isolation from words, and we know that children learn language holistically, by hearing and imitating it in complete form. Thus we can see the relationship between microtonal intervals and melody, much as we see the relationship between phonemes and words, in terms of assonance and rhyming: though in language the phonemes involved do not break our perception of words, nonetheless they provide links and points of attraction among words—which are exploited the more language becomes aesthetic rather than utilitarian.

In music—entirely aesthetic rather than utilitarian—we should expect the assonance of particular intervals to serve as a cohesive link among melodies, which is exactly what we observe: this is the dwelling in a particular *maqam* that produces one form of *tarab*, or ecstasy, in the listener. This inseparability of intonation from melody in *maqam* explains both the perceptibility of a large number of distinct intonations, as well as the simpler classification of intervals in nomenclature. In practice, I (and all of the practicing Arab musicians I’ve encountered) think about three broad categories of intonation rather than twelve or more: flat, half-flat, and natural (or natural, half-sharp, and sharp); just as in English, I think of the vowel sounds as A, E, I, O, and U, even though in speech there are many different versions of each, but I need actual words using the different versions in order to perceive or describe them. In *maqam*, I can’t play or sing each one of those many different intonations in isolation—and certainly not side-by-side—but by using the particular melodies for each *maqam* I can play each intonation correctly. And each intonation, along with the melodies using it, immediately evokes a particular *maqam*—serving as a sign for it—when I’m listening to others play or sing Arabic music. The existence of distinct melody-words I’ve stored in memory is thus an aid to the perception of intonation, as much as the distinct intonations serve as a memory aid and classification scheme for those melody-words.³

³ One caveat here: in drawing attention to the structural similarity between music and spoken language in terms of a discrete combinatorial word-based organization, I am in no way suggesting that musical words have a link to meaning in the way that each language word links to a concept or object in the world, and in no way suggesting that musical words organize into a part-of-speech-based grammar built from subjects, verbs, objects, qualities, etc. In the case of meaning in spoken language, we know already from studies in neuroscience that meaning is separately organized from word parsing and grammar: striking examples of brain damage to one processing region of the brain but not the other have resulted in cases of individuals

Variation in Words and the Fuzzy Boundaries of Language

Another parallel between music and language I'd like to explore relates to one of the **most important basic problems of cognition**: how does the mind identify something despite variations in appearance or sound, despite mistakes in part of the message, despite aging or a new haircut on a familiar face? (Pinker, 1997)⁴

Let me give you a concrete example of one dimension in which that occurs in Arabic music. Here's the melody-word I used in example 3

Example 3



And here are several ways I might vary it:

Example 10



Once a student of Arabic music has heard this phrase enough times, and with enough variations, he or she will be able, like me, to identify it in any particular incarnation. Here it is, in a variation, embedded within a longer phrase⁵:

producing grammatically correct language without any linkage to meaning whatsoever, as well as opposite cases of individuals able to communicate meaning through words but without any grammatical consistency (Pinker 1994). So even in spoken language, the discrete combinatorial organization of words is separate from the use of words as symbols representative of concepts and objects in the world—and I'm claiming that music shares one but not the other element of language cognition. In the case of part-of-speech-based grammar, there is also strong evidence that such an organization into types of words relates to the innate categorization of our perception of the world around us, into actors, actions, things, and qualities. Without such a need for music to represent the world directly, in cognition, a musical grammar would be free from such categorization; the claim I wish to develop further is that the base grammatical organization of music and language is related to a universal network-based organization of cognition generally, and that part-of-speech-based grammar is simply a local feature of the grammar of spoken language.

⁴ To those who have not encountered cognitive science, this may not seem like a problem, since it is an ordinary basic cognitive skill of humans, one we use every minute of every day. For an accessible way of understanding why this is a problem, imagine making a mistake of one character in a computer program—the computer can't recognize the intent, and produces an error. On the other hand, any normal person might not even notice the mistake at all. Cognitive and computer scientists have not yet understood how this mechanism works, although Google, with its massive parallel processing and enormous memory resources, is beginning to solve one aspect of the problem with its “did you mean...?” suggestions.

⁵ It is worth noting that the particular patterns of variation and ornamentation used are also melody-words: they function as another discrete combinatorial system built from units stored in memory by musicians and listeners. The interaction of the larger-scale melody words and the smaller-scale ornament-words is the subject for a longer discussion; this is clearly a divergence between music and spoken language, although spoken languages also have a discrete vocabulary of expressive variations for words: e.g. “puh-LEEZ!” for “please” and “It's fuh-REEZing out here.”

Example 11



Now, if I were to play the following phrase for you:

Example 12



and ask you to repeat it, as an un-experienced musician you might repeat it like this:

Example 13



This is what happens in my *maqam* classes when I give complex phrases to imitate: without getting all of the details, students are still able to catch the basic gist of the phrase. Now imagine that you repeated it to yourself over and over again without re-hearing my version. Eventually it would become a word in your musical vocabulary, identifiable and recognizable as being the same as the word in my language, but nonetheless with a slight variation. Imagine you passed it on to someone else, and he or she acquired it slightly differently: we've all played the game of "telephone"...

What do we find in spoken languages? Over geographic space, different populations pronounce certain phonemes with slight differences, some words differ from one group to another, some idiomatic phrases differ—yet people from two different regions with distinct dialects can still comprehend each other.

The same is true of orally-transmitted music traditions. The *maqam* tradition is practiced in one form or another from North Africa, through Turkey and the Levant, through Iran and Central Asia, all the way into western China. In Syria, the E-half-flat that is the third note of *maqam rast* (a note somewhere in between e-flat and e-natural) is slightly higher than the E-half-flat in *maqam rast* as played in Egypt, yet a phrase in *maqam rast* is unmistakable as such, and a Syrian will recognize an Egyptian playing *rast*, even if he also recognizes it as the Egyptian version and not his own. The differences in intonation and ornamentation are even more pronounced—while still retaining the *maqam* identity—if we compare Arabic melodies with Turkish and Greek Rebetika melodies.⁶ We can hear and identify many of the same melody-words across these traditions. In Turkey, *makam rast* uses a scale that closely resembles the Western major scale—the third scale degree, which in an Egyptian *maqam rast* is a quarter-flat, has been raised almost to natural—yet the melodies of that *maqam* are in many cases almost identical to those of *rast* in Egypt and Syria, and not to those of *maqam ajam* (the major-scale based *maqam* in the Arab world). I myself, as a practicing Arab musician, recognize Turkish *rast* melodies as *rast* melodies and not as *ajam* melodies, despite this big difference in intonation in the third note.

Another related example of the same phenomenon occurs when non-musicians sing songs in the various *maqamat*. In both Egypt and Syria I encountered many ordinary people who had an extensive knowledge of the music repertory of their culture—some could sing hundreds of songs, and it is much more common for the average person to sing than it is in the West. In Egypt,

⁶ http://media.libsyn.com/media/shumays/PPerform_011_2006-08-13.mp3 and http://media.libsyn.com/media/shumays/PPerform_012_2006-09-01.mp3 are part 1 and 2 of a Podcast I did in 2006 with Turkish *oud* player Sinan Erdemsel and a Bay Area-based Greek Rebetika/Smyrnaica group called Smyrna Time Machine, in which I compared Arabic, Greek, and Turkish versions of 3 different *maqamat*: *Rast*, *Saba*, and *Huzam/Sikah*.

everybody old enough to remember Umm Kulthum can sing parts of many of her most famous songs; many practicing Muslims can sing (recite) the call to prayer or verses of the *Qur'an* with the melodies used by the *Mu'ezziins* (those who give the call to prayer) and the *Sheikhs* who recite the *Qur'an*. In all of these cases, the *maqam* of the melody sung is perfectly clear, even when the singer's intonation is poor, and lacks the control of professionals. Even without the precision enabling a musician to use 12 different notes within the space of one semitone, the singing of amateurs manifests all of the different *maqamat* used in whatever melody they are singing, and when I or other professional musicians listen, we don't hear that singing as being in a different *maqam* when it uses a different intonation than we would use for that *maqam*—rather we hear and accept the *maqam* intended, because the melodic information conveys it, and because of the same cognitive ability that allows us to accept different pronunciations of the same word in spoken language as being the same rather than different words.⁷

My point is that we find the same variation in music as we do in language, over populations, and for the same basic reason: the cognitive ability to recognize two slightly different entities as similar enough to be identified in the same way. That variation occurs in a number of different dimensions: slightly different pronunciations of phonemes; some words that occur in the vocabulary of one dialect but not the other, and vice-versa; slightly different grammatical tendencies; etc.

I'd like to pose a fundamental question here: how can we determine the boundary of a language or of a musical practice? If we focus for the moment on vocabulary, we can see that in spoken language, different individuals from the same geographical region and around the same age and socio-economic background will likely have vocabularies that differ from each other by much less than 1%. If we move outward to people from different cities, or different ages or different backgrounds, the difference in their vocabulary will be slightly higher. Differences will be greater among people speaking different dialects of the same language. Two people speaking different languages will still have a significant portion of vocabulary in common, if those languages are related, like Spanish and Italian. And two people speaking more distantly related languages, but which have had some historical contact, will still have a portion of vocabulary in common—such as English and French. So how do we determine what is the boundary of a given language (in terms of vocabulary)?

The answer is that we cannot. The vocabulary of a language has *fuzzy boundaries*. There is a certain component, at the center, which is shared by all speakers, but as we move to the periphery, we will find words used by smaller and smaller portions of the population, until we get very peripheral words used by only a few speakers. Yet we cannot put a definitive boundary on where the center is and where the periphery is; the center for one group might be more toward the periphery for another group, and vice-versa. James Bau Graves, in his book *Cultural Democracy*, made a profound observation about cultures and populations, which applies equally well to language

⁷ It is worth noting one of the important features of language enabling the brain to treat as equivalent two slightly different things: redundancy. Redundancy adds extra padding to a signal in order to make sure that the signal is interpreted correctly. We can see how this happens easily in written text: everyone has probably received that email in which all the letters have been scrambled in each word, and it seems amazing that we can read it at all—Pinker gives the following example: “yxx cxn xndxrstxnd whxt x xm wrxtxng xvsn xf x rxplx cx xll thx vxwxls wxth xn ‘x’”—but of course in 2010 we are all used to cutting out redundant characters 2 snd txt msgs.

Ironically, it is redundancy that allows the *gradual* distortion we see in language over time and geography: because of redundancy, a mispronunciation or different pronunciation of a word won't result in misunderstanding—so the one who pronounces it differently can pass it on to others, and both can co-exist in a population.

and music. He notes that for any given cultural practice, there will be hard-core fans or practitioners, those for whom that cultural practice is a matter of daily existence. Then there are those for whom that cultural practice is a part of their lives, but not essential—in the case of music, these would be the occasional fans of a particular band, rather than the groupies and band members themselves. Then there are those at the margins, those who have just heard of the band for the first time, those who used to be fans but got over it, those who are aware of the band but never heard them, those who went to a concert once. Graves argues convincingly that to sustain that band as a cultural phenomenon, all three types of participants are necessary, and that there is a constant flux among those groups. Those from the periphery move to the center and vice-versa, and this is the way an ever-renewing dynamic audience is sustained.

We see this phenomenon in language vocabulary itself, as I've described it above. There are some words at the center, used by everybody, some towards the edges, and some at the very periphery. Yet there is a constant flux over time, as some words move from center to periphery and vice-versa. A vocabulary is *essentially dynamic*, with no fixed boundaries. We can also say the same thing about the grammar of a spoken language that we have said about vocabulary.⁸

And the same is true with phonemes, and especially vowel sounds, which differ gradually across populations. The closest parallel to musical intonation in language is vowel sounds, which emphasize different upper partials of a basic tone to produce the distinctive sound of each vowel. Just as we can calculate the exact size of a melodic interval in cents, or any other measure, we could calculate precisely the overtones emphasized by a particular vowel sound of a particular dialect. Yet there is no perfect example of a given vowel within any particular language—no mathematical ideal that vowel sounds attempt to approximate; instead the vowel sounds have very slight differences across dialects, and change very gradually within a given language over time (the most famous for English-language Linguists is the “Great Vowel Shift” which occurred between Old and Middle English). And if we were to catalogue every different pronunciation of a given vowel across different dialects of the same language, we would find tens of different pronunciations of each vowel, depending on the geographic origin of the speaker; expert linguists can identify with great precision where an individual comes from, down to the township or village level. This is exactly the same phenomenon we see in microtones in Arabic music: in the present day, the intonation of that so-called “neutral third” differs among the Egyptians, the Syrians, the Iraqis, and the Turks, within *maqam rast*, in phrases identified as melodically identical—furthermore, the intonation of that third has shifted within Egypt itself over the course of the twentieth century—descending from its more Turkish-sounding intonation at the beginning of the 20th century, to its current position.

The parallel between vowel sounds in language, and microtones in music, runs deeper: the vowel sounds of a particular dialect give important information to speakers and listeners: they strongly identify speakers from a particular region, distinguishing them from others—and not just to the ears of expert linguists, as mentioned above. The same is true of the microtones used in *maqam*-based music: by listening to the intonation of a particular musician I can identify whether he comes from Egypt, Iraq, Syria, Turkey, Iran, etc.—and just as native speakers of a language can't help but recognize almost immediately the geographic origin of other speakers of the language, especially if the accent is distinctly different, so I can't help but distinguish Syrian, Iraqi, and Egyptian musicians almost immediately by their intonation as well as by their ornamentation. That geographic information is carried constantly in every single melody-word.

⁸ linguists have observed numerous grammatical operations in use among some populations which are not in use among others speaking what might be considered the same dialect: and linguists do not claim that one version is correct, and the other wrong, but rather that each population, dialect, and idiolect has its own consistent grammar, slightly differing from others.

Intonation of Arabic scales is thus one among several parameters of the fuzzy-boundary phenomena I've described. Once we see that, we can recognize the same phenomenon occurring in music all over the world—in the microtonal scales used in some Irish music (from County Clare), some Swedish music, and in American Blues and Old Time music, to name a few I've personally encountered. To be more precise, intonation as practiced by world musicians is some kind of hybrid of Just intonation and fuzzy linguistic-style shifting, because the presence of tunable stringed instruments always insures that there will be some Just intervals in scales—and nobody would deny the strong pull that Just octaves, fourths, and fifths exert on the ear. Beyond that, the weakening strength of ever-more distant harmonic partials opens up intonational space to the kind of gradual shifting over time, space, and populations I've described here.⁹ And the fact of language and music acquisition by ear through repetition guarantees that a person studying a particular tradition long enough will use the particular regional and geographic intonations of that tradition even without being conscious of it.

Pinker discusses the importance of Saussure's concept of "the arbitrariness of the sign" for understanding language cognition (Pinker 1994: 75): as long as anybody has memorized the arbitrary sign used to signify something, a message using that sign can communicate meaning instantaneously. Music doesn't convey meaning in the same way, but it has the same arbitrariness¹⁰—as long as two people have the same musical elements in memory, a musical utterance is comprehensible. Any signal system is arbitrary—whether it is a spoken or computer language or morse code or radio signal—as long as elements can be distinguished in a discrete manner. The important factor for musical notes and intervals is *not* that they correspond to any particular acoustical phenomenon (like the intervals of Just intonation), but that they be *clearly and discretely distinguishable by the ear*—and for those who don't believe the evidence from practicing musicians, cognitive scientists have established the fact that any ordinary human ear can distinguish two tones even a few cents apart (Levitin 2006). Therefore, a wide range of intonation may serve as the basis of a musical system—including equal temperament—the ear can be trained to accept any interval used (as evidenced by the wide diversity of musical scales in the world), just as it can be trained to accept a potentially infinite range of vowel and consonant sounds for use in spoken language. We saw already that among the Arabic *maqamat* each intonation represents a different *maqam*, a whole class of melodies grouped together. Even apart from any discussion about the affective qualities of different-sized intervals, they are discretely distinguishable and used for that reason, whether that is for the purpose of evoking a particular *maqam*, a melody, an affective state, or even a familiar song or a particular musical style (in modern

⁹ One hypothesis I've come up with for the great degree of variability available for the "quarter-tone" in particular has to do with the interaction of upper partials. In a major third from C to E, the 5th partial of the C and the 4th partial of the E agree (both are E), while the 6th partial of the C (G natural) and the 5th partial of the E (G#) clash strongly. In a minor third from C to E-flat, the 5th partial of the C (E-natural) and the 4th partial of the E-flat (E-flat) clash strongly, while the 6th partial of the C and the 5th partial of the E-flat agree (both are G-natural). But for any third in between a major and minor third, there will be clash rather than agreement in *both* instances, and such double-clashing so prominent among the upper partials of both notes would probably outweigh any agreement of partials much higher up—such as between the 11th and 9th partials of the two notes. There is a wide range available to produce that strong double-clashing, which might make us consider the whole spectrum between the major and the minor third a distinct category of notes, any one of which behaves similarly from the acoustic perspective. On the other hand, harmonic agreement is not a strong consideration in Arabic music, which is for the most part mono- or hetero-phonic.

¹⁰ Music doesn't convey meaning, but it *does* convey information, in the technical sense intended by information theory: information is uncertainty or variability in a signal; any signal that has variation automatically conveys information, information that is in fact measurable. And any arbitrary set of signals can be used to do so.

Egyptian pop music, where the microtonal *maqamat* have been virtually eliminated, the occasional use of a microtonal melody immediately evokes a sense of tradition or folklore or antiqueness)—and such linkages are arbitrary in origin.

Where the possibility of diversification exists, people take advantage of it, consciously or not, whether that is in the case of pronunciation or vocabulary in language, in art, fashion, or in many other areas of human creativity. Such diversification takes place along the periphery, making use of the fuzziness of boundaries. There appears to be a gradual social and cognitive pressure to diversify intonation—as one of these available parameters—for two purposes: 1) distinguishing individuals and groups of people from one another, and 2) diversification for the purpose of system robustness, keeping multiple versions alive in order to insure the overall survival and longevity of the system. Intonation can thus be seen as one among several cognitive/evolutionary functions that exploit subtle differences in a narrow band of values for these two purposes, like physical appearance, cognitive and physical abilities, vowel sounds, idiomatic phrases, as well as cultural forms in general.

The availability of a large number of subtly different intonations is an aesthetic value as well in Arabic music, appreciated by connoisseurs, possibly because more subtle knowledge reflects a deeper immersion in the information held by only a select few. Distinct group identity is one undeniable component of aesthetic value—one that we might call snobbery. The expressive potential of different intonations each to evoke different affective states, or for a change in intonation to thwart expectations, may lead musicians to distinguish the *maqamat* from each other beyond their already-present melodic distinctness. As Scott Marcus and others observe, individual musicians even in one tradition take advantage of the possibility of using intonation in different ways, some preferring to maintain distinct intonations for each *maqam*, and some using the same intonations for multiple *maqamat*. This individual flavoring is a kind of play—partly innate and partly conscious—within the fuzzy boundaries available to Arabic music.

Oral Transmission and My Pedagogical Approach to Arabic Music

The final parallel I will draw between music and spoken language is the method of oral transmission, at which I've hinted above a number of times. Despite differences in pedagogical methods among different oral traditions, I will assert that the *cognitive* component of music acquisition is more or less identical across all world traditions of music. Whether teachers recite melodic (or rhythmic) phrases to students to repeat, or students learn without teachers simply by listening to and imitating other musicians or recordings of musicians, or students learn by doing, joining in and following along with others who know the tune already, the same thing is happening in the brains of people learning a musical language: hearing, imitation, repetition, and the development of a finer discrimination of musical elements over time. Without going into a discussion of musical grammar, it is inarguable that there is such a thing as musical syntax governing whether certain combinations of sounds work and some don't—and in an oral tradition, students learn that syntax over years of training and practice, *whether principles of syntax are explicitly stated or not*, just as in spoken language. In my own experience of learning Arabic music as an adult, as a second musical language, very few syntactical principles were made explicit for me, and melodies were simply presented whole. The melodies contained examples of correct syntax, and so by learning those whole melodies I was learning both the vocabulary of melody-words as well as the many different kinds of syntax. The few syntactical principles given to me explicitly by musicians seem in retrospect to be ridiculous oversimplifications of the musical reality—the melodies that I was taught were a much richer store of information than anything anybody ever told me.

Now, if I as a musician were simply repeating exactly the same melodies I was taught previously, there would be nothing to wonder about. But in fact, I am not repeating any of the same melodies I was taught when I improvise: I am recombining elements to produce new melodies, and I am doing so with correct musical syntax, without having been told what that is. The fact that I can absorb whole music and then produce completely new music, simply through the process of hearing, imitating, and repeating, suggests that there is an innate process going on, very similar to the process that enabled me to hear and absorb whole language as a child—without being explicitly told “this is a word” or “this is a sentence” or “this is correct syntax”—and end up producing completely new utterances in that language.

A comparison between music and language acquisition, in terms of the time required to develop fluency, illustrates my point: studies in mastery in many fields, including music, have established that 10 Years/10,000 hours is a reasonable measure of how long it takes to develop basic mastery in a given skill (Levitin: 197). By the same token, full mastery of spoken language is acquired over a period of years, and although children achieve language development at slightly different rates, it is nonetheless true that children under two can babble and say words, but not complete sentences, that children between 3 and 5 are able to speak in complete sentences with some grammatical problems and an incomplete vocabulary, that children between 8 and 11 are speaking almost exactly like adults, but that it takes another decade after that to acquire the full vocabulary and the nuances of usage of an adult (Pinker). It is a reasonable approximation to assume that acquisition time is in some sense proportional to the *amount of information* acquired in mastering a skill, and therefore related to the amount of memory usage by that skill, in terms of physical neurons and the connections among them. My point here is that music has a similar depth of complexity as does language, and must use a similar amount of memory storage, storage that must take a similar form in the brain as does spoken language.

Correct intonation is not something that is acquired immediately in musical practice, and not something that one can focus on exclusively before learning other aspects of music; a precise sense of intonation develops in parallel with the expansion of one’s musical vocabulary—this was my own experience in learning Arabic music, and it was not until I had studied it for around 5 or 6 years that my intonation became correct consistently. Part of that experience was traveling to study in Syria after having lived and studied in Egypt for a year, and being told by my teachers there that my intonation of several of the microtonal notes was too low—then returning to Egypt to find that the Egyptians found the Syrian intonation too high. That experience of comparing two close dialects of Arabic music was crucial for me, as someone learning Arabic music as a second language (as an adult rather than as a child), to develop a fine sense of intonation.

I am now nearly done with my 12th year of my experiment on myself in mastery of Arabic music, an experiment encouraged by my learning about the 10 years/10,000 hours measure around the time I started. I dismissed claims that I could never play music like an Egyptian because I wasn’t born there, and struggled to see how far I could go, carrying with me the belief that anyone from anywhere could do it with 10 years of attentive learning. I found that the process was much like learning a second spoken language—which I did with the Arabic language beginning at the same time I began to study Arabic music—and I recognized that I needed to pay special attention to certain details if I didn’t want to have an accent. The process has yielded the insights I present here—the most important of which are: 1. Intonation cannot be learned apart from melody. 2. Improvisation and composition are built from a very large vocabulary of melody-words. 3. The musical system cannot be reduced to a small set of rules, which once learned, give the key to making new music—rather, there is a large vocabulary of grammatical operations which must be learned one-by-one, like spoken language: another way of putting this is that the information content of music cannot be reduced to a simpler form, and so analysis done without actually having learned the

language itself, is inevitably oversimplified and lacks information content. 4. Learning happens without conscious effort, simply by spending time on it. 5. Learning does not happen sequentially but rather holistically, in two senses: first, one doesn't learn intonation, then melody, then improvisation—instead, one learns every part of the music a little bit at a time, and each specific skill involved contributes to the others and is in some sense inseparable from them; and second, every fragment of melody contains details about all of the aspects of music, intonation, ornamentation, melody words, syntax, as well as emotion and expressivity, so every melody has an equivalent richness and teaching power—and one can start anywhere.

I have been teaching Arabic music and *maqam* for the last several years, in addition to performing professionally, and pedagogy has been, for me, an important realm within which to test hypotheses about cognition and parallels with spoken language. My goal in teaching has been to activate the innate language-like music acquisition mechanism in students, and to prime them for Arabic music specifically. In fact, what I do as a teacher is not unique in any way, because I recognize that the methods of oral transmission that have evolved in world music traditions already take advantage of the innate ability to learn music. My classes are almost entirely call and response, with a little bit of analysis of where we are in a given *maqam*.

At the most basic level, imitation and repetition are the foundation of such learning. Imitation strengthens the ear's parsing mechanism, it adds muscle memory to aural memory, and it begins the process of memory storage.¹¹ Repetition enriches the perception of a melody through multiple imperfect renderings, allowing for the range of fuzziness we see in mature musicians and in the community of musicians as a whole. And repetition begins the transition in memory from phoneme to group of phonemes to word. The key mechanism to understand is the transition from short-term to long-term memory—from the cognitive and neurological point of view this question is important to ask, but for pedagogical purposes we must simply be aware that the transition happens and that we can affect how it happens. I see this as a teacher and a student of music: a group of units cluster together, eventually are perceived as one unit, and can then cluster with other units—a word, once perceived as such, can have other components added to it, because it functions as one unit in short-term memory.

So when I teach through imitation and repetition I am conscious of presenting the right amount of material to encourage this kind of memory acquisition of musical words. But it is important for students to hear words and phonemes in a larger context—phonemes don't make sense by themselves, the "rules" of combination can be only learned through the vocabulary of words, and details of intonation are clearer and more distinguishable in the context of melodies. Babies learn phonemes by hearing complete language, not by hearing phonemes broken down—this is one of the characteristics of the language acquisition mechanism in the brain. The same is true of words and combinations of words—the "grammar rules" of Arabic music can only be learned through the examples of word combinations. So I play phrases containing all of that detail, knowing that even when students can't imitate exactly on first listening, their brains are still absorbing the language holistically.

The other component of my teaching, equally important, is what I encourage students to do when they're not studying with me: to listen broadly and actively, to learn pieces in the repertory by ear, to study with different teachers. I try to impart to students awareness of the fact that musical

¹¹ The issue of muscle memory is important, because the physical aspect of sound production must be learned and remembered. Babies' babbling is actually muscle training for phonemes. Though we have not discussed the issue of play in this paper, we must consider the importance of play and experimentation in learning as well. Repetition is the key to instantiate the information acquired in memory. Repetition trains the muscles, and makes easy what once was hard.

language acquisition happens as the result of interaction with a large community, hearing many different versions of things, hearing and imitating the full range of variation acceptable for any given musical element.¹²

An Explanation for Arabic Intonation

I have presented an analysis of intonation that depends on the combination of several avenues of inquiry: 1) an understanding of the limitations of Just intonation and Equal Temperament in accounting for intervals, 2) an awareness of the social, geographical, and historical spread of different intervals, 3) a description of cognitive functions in language, including the operation of discrete combinatorial systems using words stored in memory, and the fuzzy-boundary phenomenon, and 4) experimental verification through pedagogy, and an understanding of the parallel functioning of oral transmission in both music and language. No one of these approaches is sufficient by itself to account for the intonation used in *maqam* scales in Arabic music.

This case study in intonation is, for me, a microcosm illustrating a cognitive linguistic approach to music: the consequence for ethnomusicology, of the view I am putting forward, is similar in many points to the discussion Adelaida Reyes gave in her essay “What Do Ethnomusicologists Do? An Old Question for a New Century.” She quotes Anthony Seeger, from his 2005 keynote lecture to SEM, advocating the idea that the ethnographic and social components of ethnomusicology cannot be considered separately from the musical and analytical components of music (Reyes 2009: 4, 11). In considering music as a “social object” (Reyes: 12), we can understand it in the same way we seek to understand language. Like language, music only exists in the parallel memory storage by members of a community, a community with members, as Graves pointed out, with differing levels of involvement in it. And just as that community has fuzzy boundaries around who belongs and who doesn’t, so too musical languages themselves have fuzzy boundaries. To be clear, I am not suggesting that music as a social object, and the structure of music, should be studied together because each study helps to clarify the other (as Seeger encouraged); rather I am suggesting that from the perspective of cognition the two are in fact *one* study—that the cognitive structure music takes is dependent on and inseparable from its use, spread (i.e. learning), and change in communities. A word, a grammatical construction, an intonation, all of these components of music are social objects (in addition to carrying social information, as we saw), and they shift over time and space for social reasons, which in turn are dependent on the way music cognition works.

This explanation as it stands is incomplete, however; in future papers I hope to elaborate on two more pieces of the puzzle that I feel are necessary to gain a fuller understanding of intonation: First, an evaluation of the information content contained in music, along the lines of information theory in mathematics, to demonstrate that it is possible to evaluate and quantify the different kinds of information carried by an audio signal *independently*; and in the case of Arabic music specifically,

¹² I was pleased to find that the Folk Music Department at the Sibelius Academy in Finland takes a similar language-based approach to developing folk musicians, as documented by Juniper Hill (2009, 86-112). This approach was inspired in part, according to Hill, by Albert Lord’s description of oral composition/improvisation “drawing from an orally transmitted store of traditional metric phrases” (ibid.: 95—this is what I refer to as vocabulary), and demonstrates an awareness of what Hill refers to as “auditory-memory-storage.” The other very important component of the curriculum is the emphasis on variants of phrases, melodies, and songs—that students need to learn many different equally valid versions of songs rather than one “static authoritative version” of a piece (96).

that an audio signal can use different semi-independent dimensions of variability and uncertainty to convey different types of information, including variability in intonation of certain notes to convey geographic, historical, and personal information about the singer or musician, while other elements of the music convey information on a the level of the discrete combinatorial system of melody-words and syntax. In other words, the *choice* of scale conveys one type of information while the *use* of the scale conveys another. The consequence of an analysis on the basis of information content is that it would force us to account for the processing of multiple kinds of information from one signal, in music cognition.

The second piece of the puzzle has to do with the networks involved in information storage and transfer. It has been shown mathematically that a very special kind of network, the “Small-World Network,” exists in many different kinds of natural systems: social relations among people (the 6-degrees-of separation phenomenon), hyperlinks on the internet, neurons in the brain, the spread of disease, and the words of the English language, among many others (Strogatz 2003). Such networks have the very powerful combination of a high degree of clustering in parts of the network (many connections among close neighbors) along with a very short average path length over the network as a whole, resulting from a few long-distance connections and the existence of hub nodes connecting many different nodes. I believe an evaluation of the networks existing in musical languages will show: 1. that melody-words are connected via a small-world network just as English words have been shown to be; 2. that the network of neurons mirrors the network of words, in the sense that the networked organization of information cannot be reduced to a simpler form when representing it; 3. that the network is built innately through the learning process that links common words to other words; 4. that the grammar and syntax of information is based on short paths within the network rather than on abstract rules, and 5. that the interaction of two networks—the external network among individuals in a community, and the internal network linking units of information within the brain of each individual—is crucial both to the robustness of this kind of information, as well as to its constant evolution (this is a more mathematical way of stating Grave’s point about the center and the periphery—hubs are the center of such networks).

In terms of intonation specifically, such an analysis can give an explanation for why it remains stable within a given population, but still changes gradually over time and geographic space; such an explanation could also account for the shifting of phoneme sounds in spoken language, and for the gradual changes in vocabulary and grammar that we actually see, changes that don’t result in incomprehensibility among speakers of slightly different dialects, or speakers of different ages. The existence of such parallel cognitive information networks in both music and spoken language could also affect how we understand learning and information evolution in general.

References

- Graves, James Bau. 2004. *Cultural Democracy: The Arts, Community, and the Public Purpose*. Champaign, IL: University of Illinois Press.
- Hill, Juniper. 2009. “Rebellious Pedagogy, Ideological Transformation, and Creative Freedom in Finnish Contemporary Folk Music.” *Ethnomusicology* 53(1):86-114.
- Lerdahl, Fred and Ray Jackendoff. 1983. *A Generative Theory of Tonal Music*. Cambridge: The MIT Press.

Levitin, Daniel J. 2006. *This is Your Brain on Music: The Science of a Human Obsession*. New York: Penguin Group, Inc.

Pierce, John R. [1961] 1980. *An Introduction to Information Theory: Symbols, Signals and Noise*. 2d ed. New York: Dover Publications, Inc.

Pinker, Steven. 1994. *The Language Instinct: How the Mind Creates Language*. New York: Perennial.

-----, 1997. *How the Mind Works*. New York: W.W. Norton & Company, Inc.

Reyes, Adelaida. 2009. "What Do Ethnomusicologists Do? An Old Question for a New Century." *Ethnomusicology* 53(1):1-17.

Shannon, Claude E. and Warren Weaver. 1949. *The Mathematical Theory of Communication*. Chicago: University of Illinois Press.

Strogatz, Steven. 2003. *Sync: The Emerging Science of Spontaneous Order*. New York: Hyperion.