

Chapter 2

Degrees of Freedom in Prosody Modeling

Yi Xu and Santitham Prom-on

Abstract Degrees of freedom (DOF) refer to the number of free parameters in a model that need to be independently controlled to generate the intended output. In this chapter, we discuss how DOF is a critical issue not only for computational modeling, but also for theoretical understanding of prosody. The relevance of DOF is examined from the perspective of the motor control of articulatory movements, the acquisition of speech production skills, and the communicative functions conveyed by prosody. In particular, we explore the issue of DOF in the temporal aspect of speech and show that, due to certain fundamental constraints in the execution of motor movements, there is likely minimal DOF in the relative timing of prosodic and segmental events at the level of articulatory control.

2.1 Introduction

The ability to model speech prosody with high accuracy has long been the dream of prosody research, both for practical applications such as speech synthesis and recognition and for theoretical understanding of prosody. A key issue in prosody modeling is degrees of freedom (henceforth interchangeable with DOF). DOF refers to the number of independent parameters in a model that needs to be estimated in order to generate the intended output. So far there has been little serious discussion of the issue of DOF in prosody modeling, especially in terms of its theoretical implications. Nevertheless, DOF is often implicitly considered, and it is generally believed that, other things being equal, the fewer degrees of freedom in a model the better. For example, in the framework of intonational phonology, also known as the AM theory or the Pierrehumbert model of intonation, it is assumed that, at least for nontonal languages like English, “sparse tonal specification is the key to combining accurate phonetic modeling with the expression of linguistic equivalence

Y. Xu (✉)

University College London, London, UK

e-mail: yi.xu@ucl.ac.uk

S. Prom-on

King Mongkut's University of Technology, Thonburi, Thailand

© Springer-Verlag Berlin Heidelberg 2015

K. Hirose, J. Tao (eds.), *Speech Prosody in Speech Synthesis: Modeling and generation of prosody for high quality and flexible speech synthesis*, Prosody, Phonology and Phonetics, DOI 10.1007/978-3-662-45258-5_2

of intonation contours of markedly different lengths” (Arvaniti and Ladd 2009, p. 48). The implication of such sparse tonal representation is that there is no need to directly associate F_0 events with individual syllables or words, and for specifying F_0 contours between the sparsely distributed tones. This would mean high economy of representation. Sparse F_0 specifications are also assumed in various computational models (e.g., Fujisaki 1983; Taylor 2000; Hirst 2005).

A main feature of the sparse tonal specification is that prosodic representations are assigned directly to surface F_0 events such as peaks, valleys, elbows, and plateaus. As a result, each temporal location is assigned a single prosodic representation, and an entire utterance is assigned a single string of representations. This seems to be representationally highly economical, but it also means that factors that do not directly contribute to the major F_0 events may be left out, thus potentially missing certain critical degrees of freedom. Another consequence of sparse tonal specification is that, because major F_0 events do not need to be directly affiliated with specific syllables or even words, their timing relative to the segmental events has to be specified in modeling, and this means that temporal alignment constitutes one or more (depending on whether a single point or both onset and offset of the F_0 event need to be specified) degrees of freedom. Thus many trade-offs need to be considered when it comes to determining DOF in modeling.

In this chapter we take a systematic, though brief look at DOF in prosody modeling. We will demonstrate that DOF is not only a matter of how surface prosodic events can be economically represented. Rather, decisions on DOF of a model should be ecologically valid, i.e., with an eye on what human speakers do. We advocate for the position that every degree of freedom (DOF) needs to be independently justified rather than based only on adequacy of curve fitting. In the following discussion we will examine DOF from three critical aspects of speech: motor control of articulatory movements, acquisition of speech production skills, and communicative functions conveyed by prosody.

2.2 The Articulatory Perspective

Prosody, just like segments, is articulatorily generated, and the articulatory process imposes various constraints on the production of prosodic patterns. These constraints inevitably introduce complexity into surface prosody. As a result, if not properly understood, the surface prosodic patterns due to articulatory constraints may either unnecessarily increase the modeling DOF or hide important DOF. Take F_0 for example. We know that both local contours and global shapes of intonation are carried by voiced consonants and vowels. Because F_0 is frequently in movement, either up or down, the F_0 trajectory within a segment is often rising or falling, or with an even more complex shape. A critical question from an articulatory perspective is, how does a voiced segment get its F_0 trajectory with all the fine details? One possibility is that all F_0 contours are generated separately from the segmental string of speech, as assumed in many models and theories, either explicitly or implicitly, and especially

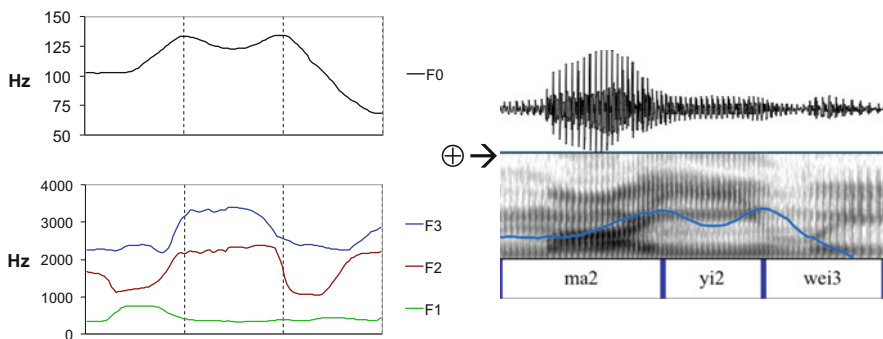


Fig. 2.1 Left: Continuous F_0 (top) and formant (bottom) tracks of the Mandarin utterance “(bi3) ma2 yi2 wei3 (shan4)” [More hypocritical than Aunt Ma]. Right: Waveform, spectrogram and F_0 track of the same utterance. Raw data from Xu (2007)

in those that assume sparse tonal specifications (Pierrehumbert 1980; Taylor 2000; ‘t Hart et al. 1990). This scenario is illustrated in Fig. 2.1, where continuous F_0 and formant contours of a trisyllabic sequence are first separately generated with all the trajectory details (1a), and then merged together to form the final acoustic output consisting of both formant and F_0 trajectories. The critical yet rarely asked question is, is such an *articulate-and-merge* process biomechanically possible?

As is found in a number of studies, the change of F_0 takes a significant amount of time even if the speaker has used maximum speed of pitch change (Sundberg 1979; Xu and Sun 2002). As found in Xu and Sun (2002), it takes an average speaker around 100 ms to make an F_0 movement of even the smallest magnitude. In Fig. 2.1, for example, the seemingly slow F_0 fall in the first half of syllable 2 is largely due to a necessary transition from the high offset F_0 due to the preceding rising tone to the required low F_0 onset of the current rising tone, and such movements are likely executed at maximum speed of pitch change (Kuo et al. 2007; Xu and Sun 2002). In other words, these transitional movements are mainly due to articulatory inertia. Likewise, there is also evidence that many of the formant transitions in the bottom left panel of Fig. 2.1 are also due to articulatory inertia (Cheng and Xu 2013).

Given that F_0 and formant transitions are mostly due to inertia, and are therefore by-products of a biomechanical system, if the control signals (from the central nervous system (CNS)) sent to this system also contained all the inertia-based transitions, as shown on the left of Fig. 2.1, *the effect of inertia would be applied twice*. This consideration makes the *articulate-and-merge* account of speech production highly improbable. That is, it is unlikely that continuous surface F_0 contours are generated (with either a dense or sparse tonal specification) independently of the segmental events, and are then added to the segmental string during articulation.

But how, then, can F_0 contours and segmental strings be articulated together? One hypothesis, as proposed by Xu and Liu (2006), is that they are coproduced under the coordination of the syllable. That is, at the control level, each syllable is

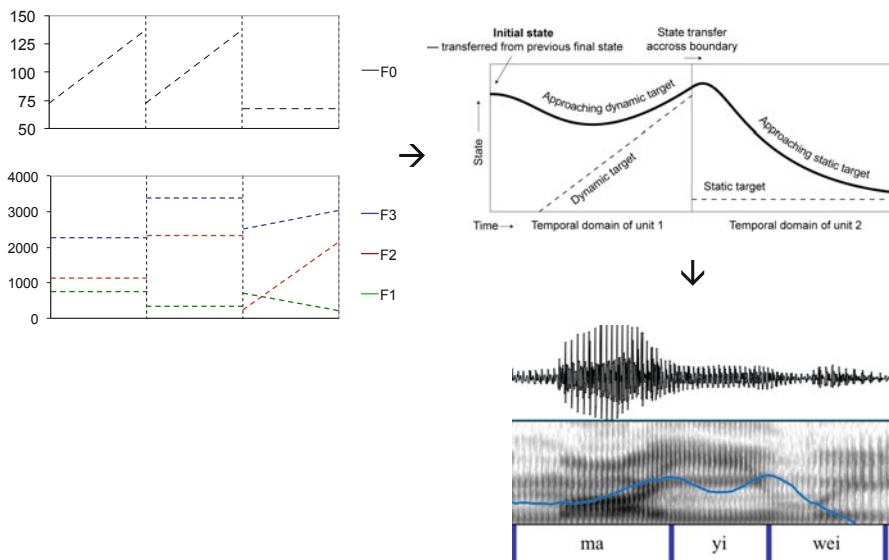


Fig. 2.2 *Left*: Hypothetical underlying pitch (*upper*) and formant (*lower*) targets for the Mandarin utterance shown in Fig. 2.1. *Top right*: The target approximation (TA) model (Xu and Wang 2001). *Bottom right*: Waveform, spectrogram and F_0 track of the same utterance. Raw data from Xu (2007)

specified with all the *underlying articulatory targets* associated with it, including segmental targets, pitch targets, and even phonation (i.e., voice quality) targets. This is illustrated in the top left block of Fig. 2.2 for pitch and formants. Here the formant patterns are representations of the corresponding vocal tract shapes, which are presumably the actual targets. The articulation process then concurrently approaches all the targets, respectively, through target approximation (top right). The target approximation process ultimately generates the continuous surface F_0 and formant trajectories (bottom right), which consist of mostly transitions toward the respective targets. Thus, every syllable, before its articulation, would have been assigned both segmental and suprasegmental targets as control signals for the articulatory system. And importantly, the effects of inertia are applied only once, during the final stage of articulatory execution.

Pitch target specification for each and every syllable may mean greater DOF than the sparse tonal specification models, of course, which is probably one of the reasons why it is not widely adopted in prosody modeling. But what may not have been apparent is that it actually reduces a particular type of DOF, namely, the F_0 -segment alignment. For the sparse tonal specification models, because F_0 events are not attached to segments or syllables, the relative alignment of the two becomes a free variable, which constitutes at least one DOF (two if onset and offset of an F_0 event both have to be specified, as in the Fujisaki model). Thus for each tonal event, not only its height, but also its position relative to a segmental event, need to be specified. This complexity is further increased by the assumption of most of the sparse-tonal

specification models that the number of tonal and phrasal units is also a free variable and has to be learned or specified. For the Fujisaki model, for example, either human judgments have to be made based on visual inspection (Fujisaki et al. 2005), or filters of different frequencies are applied first to separately determine the number of phrase and accent commands, respectively (Mixdorff 2000). Even for cases where pitch specifications are obligatory for each syllable, e.g., in a tone language, there is a further question of whether there is freedom of micro adjustments of F_0 -segment alignment. Allowance for micro alignment adjustments is assumed in a number of tonal models (Gao 2009; Gu et al. 2007; Shih 1986).

There has been accumulating evidence against free temporal alignment, however. The first line of evidence is the lack of micro alignment adjustment in the production of lexical tones. That is, the unidirectional F_0 movement toward each canonical tonal form starts at the syllable onset and ends at syllable offset (Xu 1999). Also the F_0 -syllable alignment is not affected by whether the syllable has a coda consonant (Xu 1998) or whether the syllable-initial consonant is voiced or voiceless (Xu and Xu 2003). Furthermore, the F_0 -syllable alignment does not change under time pressure, even if tonal undershoot occurs as a result (Xu 2004). The second line of evidence is from motor control research. A strong tendency has been found for related motor movements to be synchronized with each other, especially when the execution is at a high speed. This is observed in studies of finger tapping, finger oscillation, or even leg swinging by two people monitoring each other's movements (Kelso et al. 1981; Kelso 1984; Kelso et al. 1979; Mechsner et al. 2001; Schmidt et al. 1990). Even non-cyclic simple reaching movements conducted together are found to be fully synchronized with each other (Kelso et al. 1979).

The synchrony constraints could be further related to a general problem in motor control. That is, the high dimensionality of the human motor system (which is in fact true of animal motor systems in general) makes the control of any motor action extremely challenging, and this has been considered as one of the central problems in the motor control literature (Bernstein 1967; Latash 2012). An influential hypothesis is that the CNS is capable of functionally freezing degrees of freedom to simplify the task of motor control as well as motor learning (Bernstein 1967). The freezing of DOF is analogous to allowing the wheels of a car to rotate only around certain shared axes, under the control of a single steering wheel. Thus the movements of the wheels are fully synchronized, and their degrees of freedom merged. Note that such synchronization also freezes the relative timing of the related movements, hence eliminating it as a DOF. This suggests that the strong synchrony tendency found in many studies (Kelso et al. 1979; Kelso et al. 1981; Mechsner et al. 2001; Schmidt et al. 1990) could have been due to the huge benefits brought by the reduction of temporal degrees of freedom.

The benefit of reducing temporal DOF could also account for the tone-syllable synchrony in speech found in the studies discussed above. Since articulatory approximations of tonal and segmental targets are separate movements that need to be produced together, they are likely to be forced to synchronize with each other, just like in the cases of concurrent nonspeech motor movements. In fact, it is possible that

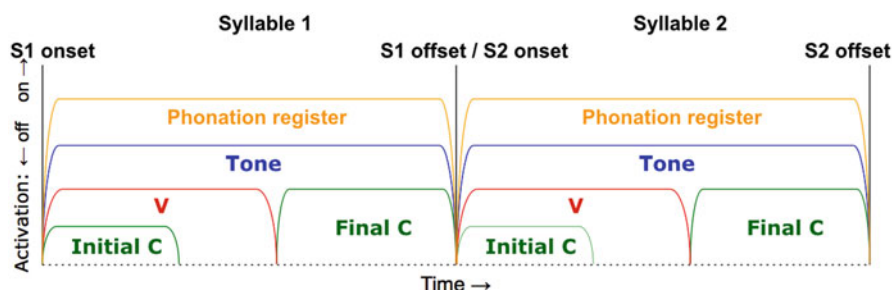


Fig. 2.3 The time structure model of the syllable (Xu and Liu 2006). The syllable is a *time structure* that assigns temporal intervals to consonants, vowels, tones, and phonation registers (each constituting a phone). The alignment of the temporal intervals follows three principles: **a** *Co-onset* of the initial consonant, the first vowel, the tone, and the phonation register at the beginning of the syllable; **b** *Sequential offset* of all noninitial segments, especially coda C; and **c** *Synchrony* of laryngeal units (tone and phonation register) with the entire syllable. In each case, the temporal interval of a phone is defined as the time period during which its target is approached

the syllable is a mechanism that has evolved to achieve synchrony of multiple articulatory activities, including segmental, tonal, and phonational target approximations. As hypothesized by the *time structure model of the syllable* (Xu and Liu 2006), the syllable is a temporal structure that controls the timing of all its components, including consonant, vowel, tone, and phonation register (Xu and Liu 2006), as shown in Fig. 2.3. The model posits that the production of each of these components is to articulatorily approach its ideal target, and the beginning of the syllable is the onset of the target approximation movements of most of the syllabic components, including the initial consonant, the first vowel, the lexical tone, and the phonation register (for languages that use it lexically). Likewise, the end of the syllable is the offset of all the remaining movements. In this model, therefore, there is always full synchrony at the onset and offset of the syllable. Within the syllable, there may be free timing at two places, the offset of the initial consonant, and the boundary between the nuclear vowel and the coda consonant. In the case of lexical tone, it is also possible to have two tonal targets within one syllable, as in the case of the L tone in Mandarin, which may consist of two consecutive targets when said in isolation. The boundary between the two targets is probably partially free, as it does not affect synchrony at the syllable edges.

2.3 The Learning Perspective

Despite the difficulty of motor control just discussed, a human child is able to acquire the ability to speak in the first few years of life, without formal instructions, and without direct observation of the articulators of skilled speakers other than the visible ones such as the jaw and lips. The only sure input that can inform the child of the articulatory details is the acoustics of speech utterances. How, then, can the child

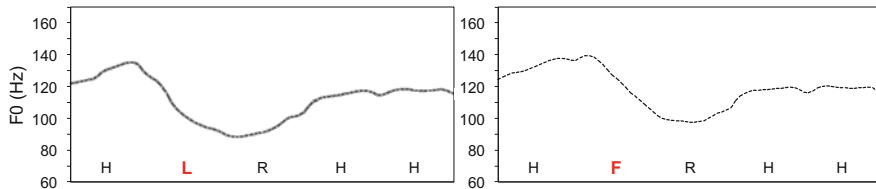


Fig. 2.4 Mean time-normalized F_0 contours of five-syllable Mandarin sentences with syllable two carrying the low tone on the *left* and the falling tone on the *right*. Data from Xu (1999)

learn to control her own articulators to produce speech in largely the same way as the model speakers? One possibility is that the acquisition is done through analysis-by-synthesis with acoustics as the input. The strategy is also known as distal learning (Jordon and Rumelhart 1992). To be able to do it, however, the child has to face the problem of multiplicity of DOF. As discussed earlier, adult speech contains extensive regions of transitions due to inertia. Given this fact, how can the child know which parts of the surface contour are mostly transitions, and which parts best reflect the targets? In Fig. 2.4, for example, how can a child tell that the Mandarin utterance on the left contains a low tone, while the one on the right contains a falling tone at roughly the same location? One solution for the child is to confine the exploration of each tonal target to the temporal domain of the syllable. That way, the task of finding the underlying target is relatively simple. This strategy is implemented in our computational modeling of tone as well as intonation (Liu et al. 2013; Prom-on et al. 2009; Xu and Prom-on 2014). Our general finding is that, when the syllable is used as the tone-learning domain, their underlying targets are easily and accurately extracted computationally, judging from the quality of synthesis with the extracted tonal targets in all these studies.

The ease of extracting tonal targets within the confine of the syllable, however, does not necessarily mean that it is the best strategy. In particular, what if the synchronization assumption is relaxed so that the learning process is given some freedom in finding the optimal target-syllable alignment? In the following we will report the results of a modeling experiment on the effect of flexibility of timing in pitch target learning.

2.3.1 *Effect of Freedom of Tone–Syllable Alignment on Target Extraction—An Experiment*

The goal of this experiment is to test if relaxing strict target-syllable synchrony improves or reduces F_0 modeling accuracy and efficiency with an articulatory-based model. If there is real timing freedom either in production or in learning, modeling



Fig. 2.5 Illustration of *onset timing shifts* used in the experiment and their impacts on the timing of adjacent syllables

accuracy should improve with increased timing flexibility during training. Also assuming that there is regularity in the target alignment in mature adults' production, the process should be able to learn the alignment pattern if given the opportunity.

2.3.1.1 Method

To allow for flexibility in target alignment, a revised version of PENTAtainer1 (Xu and Prom-on 2010–2014) was written. The amount of timing freedom allowed was limited, however, as shown in Fig. 2.5. Only onset alignment relative to the original was made flexible. For each syllable, the onset of a pitch target is either set to be always at the syllable onset (fixed alignment), or given a 50 or 100 ms search range (flexible alignment). In the case of flexible alignments, if the hypothetical onset is earlier than the syllable onset, as shown in row two, the synthetic target approximation domain becomes longer than that of the syllable, and the preceding target domain is shortened; if the hypothetical onset is later than the syllable onset, as shown in row three, the synthetic target approximation domain is shortened, and the preceding domain is lengthened. Other, more complex adjustment patterns would also be possible, of course, but even from this simple design, we can already see that the adjustment of any particular alignment has various implications not only for the current target domain, but also for adjacent ones.

The training data are from Xu (1999), which have been used in Prom-on et al. (2009). The dataset consists of 3840 five-syllable utterances recorded by four male and four female Mandarin speakers. In each utterance, the first two and last two syllables are disyllabic words while the third syllable is a monosyllabic word. The first and last syllables in each sentence always have the H tone while the tones of the other syllables vary depending on the position: H, R, L, or F in the second syllable, H, R, or F in the third syllable, and H or L in the fourth syllable. In addition to tonal variations, each sentence has four focus conditions: no focus, initial focus, medial focus, and final focus. Thus, there are 96 total variations in tone and focus.

2.3.1.2 Results

Figure 2.6 displays bar graphs of RMSE and correlation values of resynthesis performed by the modified PENTAtainer1 using the three onset time ranges. As can be seen, the 0 ms condition produced lower RMSE and higher correlation than

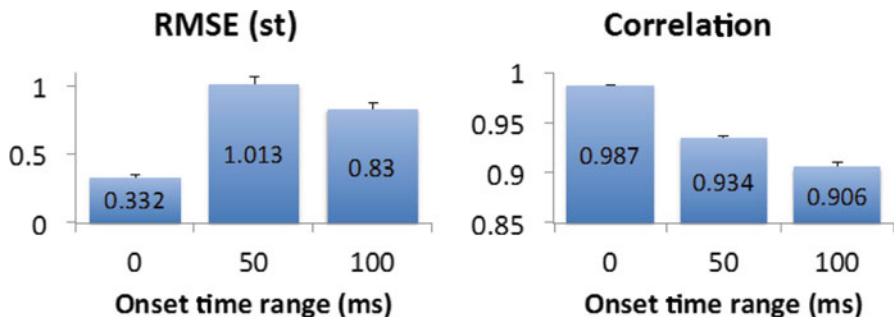


Fig. 2.6 Root mean square error (*RMSE*) and Pearson’s correlation in resynthesis of F_0 contours of Mandarin tones in connected speech (data from Xu 1999) using targets obtained with three onset time ranges: 0, 50, and 100 ms

both the 50 and 100 ms conditions. Two-way repeated measures ANOVAs showed a highly significant effect of onset time range on both RMSE ($F [2, 7] = 387.4$, $p < 0.0001$) and correlation ($F [2, 7] = 320.1$, $p < 0.0001$). Bonferroni/Dunn post-hoc tests showed significant differences between all onset time ranges for both RMSE and correlation. More interestingly, on average, the learned alignments in the flexible conditions are still centered around syllable onset. The average deviation from the actual syllable boundaries is -2.3 ms in the 50 ms onset range condition and -5.1 ms in the 100 ms onset range condition (where the negative values mean that the optimized onset is earlier than the syllable boundary). A similarly close alignment to the early part of the syllable has also been found in Cantonese for the accent commands in the Fujisaki model, despite lack of modeling restrictions on the command–syllable alignment (Gu et al. 2007).

Figure 2.7 shows an example of curve fitting with 0 and 100 ms onset shift ranges. As can be seen, pitch targets learned with the 0 onset shift range produced a much tighter curve fitting than those learned with free timing. More importantly, we can see why the increased onset time range created problems. For the third syllable, the learned optimal onset alignment is later than the syllable onset. As a result, the temporal interval for realizing the preceding target is increased, given the alignment adjustment scheme shown in Fig. 2.6. As a result, the original optimal F_0 offset is no longer optimal, which leads to the sizeable discrepancy in Fig. 2.7b. Note that it is possible for us to modify the learning algorithm so that once an optimized onset alignment deviates from the syllable boundary, the preceding target is reoptimized. However, that would lead to a number of other issues. Should this reoptimization use fixed or flexible target onset? If the latter, shouldn’t the further preceding target also be reoptimized? If so, the cycles will never end. Note also, that having a flexible search range at each target onset already increases the number of searches by many folds; having to reapply such searches to earlier targets would mean many more folds of increase. Most importantly, these issues are highly critical not just for modeling, but also for human learners, because they, too, have to find the optimal targets during their vocal learning.

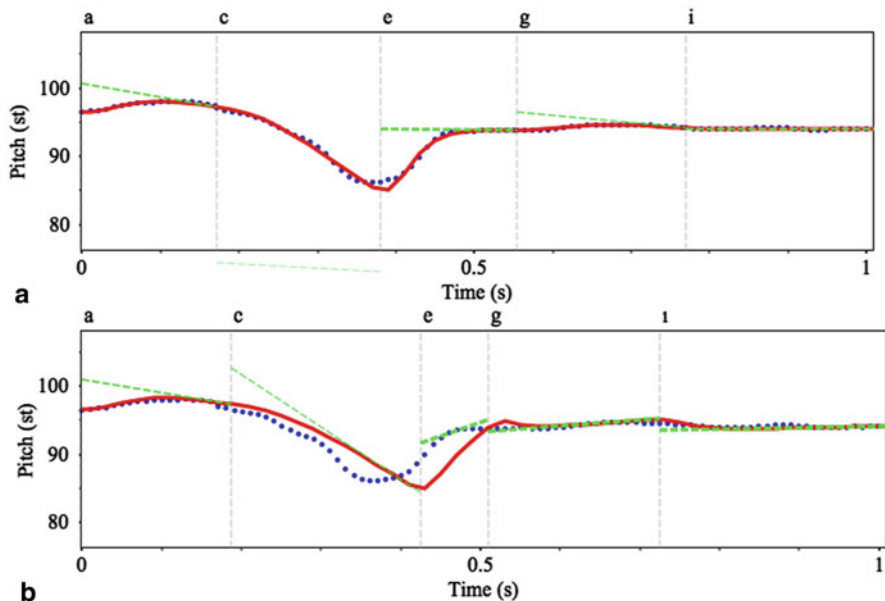


Fig. 2.7 Examples of curve fitting with targets learned with 0 ms onset timing shift (a), and 100 ms shift (b). The *blue dotted lines* are the original contours and the *red solid lines* the synthetic ones

In summary, the results of this simple modeling experiment demonstrate the benefit of fixing the temporal domain of the tonal target to that of the syllable in tone learning. Fully synchronizing the two temporal domains reduces DOF, simplifies the learning task, shortens the learning time, and also produces better learning results.

2.4 The Functional Perspective

Given that prosody is a means to convey communicative meanings (Bailly and Holm 2005; Bolinger 1989; Hirst 2005), the free parameters in a prosody model should be determined not only by knowledge of articulatory mechanisms, but also by consideration of the communicative functions that need to be encoded. Empirical research so far has established many functions that are conveyed by prosody, including lexical contrast, focus, sentence type (statement versus question), turn taking, boundary marking, etc. (Xu 2011). Each of these functions, therefore, needs to be encoded by specific parameters, and all these parameters would constitute separate degrees of freedom. In this respect, a long-standing debate over whether prosody models should be linear or superpositional is highly relevant. The linear approach, as represented by the autosegmental-metrical (AM) theory (Ladd 2008; Pierrehumbert 1980; Pierrehumbert and Beckman 1988), is based on the observation that speech intonations manifest clearly visible F_0 peaks, valleys, and plateaus. It is therefore assumed that

Speech Prosody in Speech Synthesis: Modeling and
generation of prosody for high quality and flexible
speech synthesis

Hirose, K.; Tao, J. (Eds.)

2015, VIII, 213 p. 60 illus. in color., Hardcover

ISBN: 978-3-662-45257-8