



SYNTHESIS: MODELLING VARIABILITY AND CONSTRAINTS

Rolf Carlson

Department of Speech Communication and Music Acoustics
Royal Institute of Technology (KTH) Stockholm, Sweden
Currently at: Spoken Language Systems Group
Laboratory for Computer Science, MIT, Cambridge, MA, USA

ABSTRACT

This paper discusses some important issues in current speech synthesis research. Modelling of speaker characteristics and emotions are used as examples of new trends in the area. The relation to speech recognition research is also emphasized. New methods such as automatic learning and the use of new analysis techniques are also referred to.

INTRODUCTION

The title of this paper might at first glance seem to be a mistake. Concepts such as variability and constraints are traditionally more related to speech recognition than speech synthesis. Variability is something that creates problems in speech recognition and many different methods have been developed to process speech in such a way that the variability to some extent can be handled. Similarly, constraint is an often-used term in speech recognition. Speech synthesis research has only recently started to deal with these two concepts.

It is a basic goal for speech research to understand when variability is allowed and when constraints are applied. It is also an important task for speech synthesis development to model the cause of variation: Is it a free variation or is it the result of a specific circumstance? Constraints have many different shapes. The freedom and the limitation of the vocal tract shape has been studied for many years. The constraints can in this case be expressed by size and mass limitations. Other useful constraints can be in the form of possible control parameter combinations in a formant type synthesizer. The move to explore higher level parameters [49] is an example of how constraints are introduced into the control structure itself rather than by explicitly formulated rules. The description of prosody in terms of synchronized and unsynchronized turning points is another example of how constraints described in autosegmental phonology terminology has had a positive influence on the speech synthesis development [10,25,34,40].

Modelling of variability is a new trend in speech synthe-

sis. Speaker characteristics are beginning to play a more important role in the specification of a text-to-speech system. Similarly inter-speaker variation is put into focus as a way of improving the naturalness of the synthesis. Emphasis, focus and emotions are starting to be important concepts in the speech synthesis community. Better understanding of these areas will have an impact on several applications in speech technology in terms of improved quality. A systematic account of speech variability helps in creating speaker adaptable speech recognition systems and more flexible synthesis schemes.

In this paper we will give some examples from the speech synthesis area where this type of thinking has been productive. It is clear that speech synthesis research has changed during the last ten years. After rather slow progress we now have a very productive phase with many new directions. The special workshops in Autrans¹ and Edinburgh² gave many good examples of this new trend. The special issue of Journal of Phonetics³ is another manifestation of the current interest in speech synthesis research.

SYNTHESIZERS AND CONTROL PARAMETERS

We currently have a number of different classes of synthesizers in our systems. The long term goal of having articulatory synthesis is also attracting considerable research effort. It is generally agreed upon that the ultimate goal, to use an articulatory-based synthesizer, will be the best solution. However we are still far from including these models into our text-to-speech systems. One problem for this development is still the lack of articulatory data. Despite new analysis methods, the data collection is one of many bottlenecks. The efforts to use neural networks to go directly

¹The ESCA workshop on Speech Synthesis in Autrans (France) September 1990

²The ESCA workshop on Speaker Characterization in Speech Technology in Edinburgh (UK) June 1990

³Journal of Phonetics Special Issue: Speech Synthesis and Phonetics, Volume 19 No 1 January 1991

from the speech waveform to articulatory parameters are thus of considerable interest [3,41]. It is clear that we will see many papers of this nature in the future.

In the other end of the synthesizer continua we have the PSOLA type of method [17]. The algorithms are based on a pitch-synchronous overlap-add approach for modifying the speech prosody and concatenating diphone waveforms. The frequency domain approach is used to modify the spectral characteristics of the signal while the time domain approach provides efficient solutions for real time implementation of synthesis systems. The PSOLA method has been very successfully applied in high quality text-to-speech synthesis systems [36]. However, there are limitations in these approaches. Speaker transformation and unit selection can cause serious problems.

Despite the difficulties in controlling formant-based synthesizers they are still used by many researchers. For example such synthesizers are used in the ESPRIT-project polyglot [10] and in other multilingual efforts [16,28]. The current formant-based synthesizer systems are slowly incorporating some of the regularities found in true articulation. This is especially the case with the glottal source models [14,21,32,50].

Since the control of a formant synthesizer can be a very complex task, some efforts have been made to help the users. The introduction of "higher level parameters" should be mentioned in this context [49]. These parameters can be used at an intermediate level that is more understandable from the user's point of view compared to the detailed synthesizer specifications. Thus, the first goal is to find a framework to simplify the process and to incorporate within the synthesis process the constraints that are known to exist. A formant frequency should not have to be adjusted specifically by the rule developer depending on nasality or glottal opening. This type of adjustment might be better handled automatically according to a well-specified model. The same process should occur with other parameters such as bandwidths and glottal settings.

The second goal for the introduction of higher level parameters is more basic in terms of understanding of the relation between the two levels of controls. This requires detailed understanding of the relation between acoustic and articulatory phonetics.

As a small test of this type of articulatory-based thinking, test stimuli along different feature dimensions have been synthesized [4]. One intention was to illustrate the power of higher level parameters in speech synthesis. It was hypothesized that the higher level parameters explored more natural dimensions than the lower level controls. Thus, the phoneme identification of intermediate stimuli should be easier for the subjects in these experiments compared to similar experiments carried out before. It was concluded

that the transition between two phoneme identities along the continuum was very abrupt, supporting this view.

MODELLING OF THE GLOTTAL SOURCE

During the last decade we have seen a strong effort to study the glottal source. In addition to the Vocal Fold Symposiums, special sections dealing with this subject have been arranged at ASA meetings and also at the Spoken Language Processing meeting in Kobe, Japan. It has been felt that understanding of the glottal source is one of the most important goals in speech synthesis work. The unnatural quality of the synthetic speech has to a large extent been blamed on a simplified glottal source. The work to synthesize female voices has supported this view [15,30,31]. Several glottal models have been proposed [21,32]. The improvement of speech quality by including an elaborate glottal model has in some cases been very impressive. It is clear that the simple models used up to now have been a major obstacle and that source modelling will be a critical aspect in the next generation of synthesis systems.

However, despite the current emphasis on source modelling, it should be noted that other aspects have equal importance. For example, improved models of the higher vocal tract resonances or the fricative spectrum in formant-type synthesis have a very strong impact on the speech quality.

It should be emphasized that quality improvement can be made in many different ways. Correct intonation can in some cases lead to high acceptability of the synthetic speech despite segmental problems. On the other hand, inferior quality in synthesis with many unnatural discontinuities and missing cues can not be "hidden" by a good prosodic model.

SPEAKER CHARACTERISTICS

Synthesis research has, to some extent, changed direction during recent years. The emphasis on CV syllables has been reduced and general aspects such as speaker characteristics, prosodic models and linguistic analysis have been given higher priority. The reasons for this change are many. One obvious reason is the limited success in enhancing the general speech quality by only improving the segmental models. The speaker-specific aspects are regarded as playing a very important role in the acceptability of synthetic speech. This is especially true when the systems are used to signal semantic and pragmatic knowledge.

One interesting effort to include speaker characteristics in a complex system has been reported by the ATR group in Japan. The basic concept is to preserve speaker characteristics in interpreting systems [1]. The proposed voice conversion technique consists of two steps: mapping codebook

generation of LPC parameters and a conversion synthesis using the mapping code book. The effort has stimulated much discussion, especially considering the application as such. The method has been extended from a frame-by-frame transformation to a segment-by-segment transformation [2].

One concern with this type of effort is that the speaker characteristics specified through training without any specific underlying model of the speaker. It would be helpful if the speaker characteristics could be modeled by a limited number of parameters. Only a small number of sentences might in this case be needed to adjust the synthesis to one specific speaker. Thus, it is a challenge for the future to find the best way to classify a speaker. The needs in both speech synthesis and speech recognition are very similar in this respect.

Several studies have recently been published concerning how a speaker adjusts to the listener and to the environment. A speaker is expected to vary the speech along a continuum of hypo- and hyperspeech [35]. It is argued that one important research task is to study the sufficient discriminability needed for communication rather than the notion of phonetic invariance. Duration-dependent vowel reduction has been one topic of research in this context. However, it seems that vowel reduction as a function of speech tempo is a speaker-dependent factor [22,47].

Duration and intonation structures and pause insertion strategies reflecting variability in the dynamic speaking style are other important speaker dependent factors [20,43,48]. Parameters like consonant-vowel ratio and source dynamics are typical parameters that have to be considered in addition to basic physiological variation. The ultimate test of our descriptions is our ability to successfully synthesize not only different voices but also different styles [5]. Appropriate modelling of these factors will increase both naturalness and intelligibility of a synthetic speech.

SYNTHESIS OF EMOTIONS

In acoustic-phonetic research most studies deal with function and realization of linguistic elements. With a few exceptions, (e.g., [46,52],) the acoustics of emotions have not been extensively studied. Most studies have dealt with the task of identifying extralinguistic dimensions qualitatively. Sometimes these studies have also included efforts to quantify these dimensions, by using scaling methods for example. Spontaneous speech has been used as well as read speech with simulated emotional expressions in these experiments.

An interesting alternative is to ask subjects to adjust test stimuli to some internal reference, such as joy, anger etc. This is typically done by using synthetic speech. The speech should not be of too poor quality if emotions should be conveyed. Recent experiments using DECtalk have been

reported by Cahn [11]. The special "affect-editor" was developed to control the synthesizer. Its success in generating recognizable affect was confirmed in an experiment in which the affect intended was perceived as such for the majority of the presentations.

The amount of interaction between the emotive speech and the linguistic content of a sentence is difficult to ascertain, but has to be taken into account. The voice does not always give away the speaker attitude. It is often observed that misinterpretation of emotions occurs if the listener is perceiving the speech signal without reference to visual cues. Depending on contextual references, it is thus easy to confuse anger with joy, fright with sorrow, etc.

Systematic variation in speech synthesis has been used as a tool to explore possible style and speaker dimensions [23]. Preliminary listening experiments were carried out with the aim of describing different synthesis samples according to different attitudinal and emotional dimensions. It was shown that such a method can be extremely valuable to explore extralinguistic types of variations.

AUTOMATIC LEARNING

We have recently noticed very interesting efforts to collect segmental data for synthesis with the help of automatic procedures. Formant-type synthesis has traditionally been based on very labor-intensive optimization work. The notion "analysis by synthesis" has not been explored except by manual comparisons between hand-tuned spectral slices and a reference spectra. The work by Holmes [26] is a good example of how to speed up this process. With the help of a synthesis model, the spectra is automatically matched against analyzed speech. The matching is done on a linear power scale to emphasize the importance of spectral peaks. The ambition is to make a broad collection of such analyzed segments and to use a clustering technique to reduce the size of the collection. Automatic techniques such as this will probably also play an important role in making speaker-dependent adjustments. One advantage with these methods is that the optimization is done in the same framework as that to be used in the production. The synthesizer constraints are thus already imposed in the initial state.

Methods for pitch-synchronous analysis will be of major importance in this context. Experiments such as the one presented by Talkin and Rowley [51] will lead to better estimates of pitch and vocal tract shape. These automatic procedures will, in the future make it possible to gather a large amount of data. Lack of glottal source data is currently a major obstacle for the development of speech synthesis with improved naturalness.

Given that we have a collection of parameter data from an analyzed speech corpora, we are in a good position to look for coarticulation rules and context-dependent varia-

tions. Detailed analysis work such as the study of vowels by Huang [27] can be complemented with automatic procedures. Rule extraction algorithms such as the one described by Bosch [9] can be applied to these types of data.

The collection of huge speech corpora has also facilitated a new possibility to test duration and intonation models on a grand scale [13,29,42,45]. Some of the old "knowledge" has been revised in this context. The new type of methods can easily create large amounts of analysis results. It will be the task for the speech synthesis researcher to summarize these in understandable models that can be used in the next generation of synthesizers [12,18,33].

UNIT SIZE

A special method to generate an allophone inventory has been proposed by the research group at NTT in Japan [24,37]. The synthesis allophones are selected with the help of the context-oriented clustering method, COC. The COC searches for the phoneme sequences of different sizes that most affect the phoneme realization. The system developed using these synthesis units was regarded to have superior speech quality compared to an earlier synthesis system based on diphones.

The context-oriented clustering approach is a good illustration of a new trend in speech synthesis. Our studies are concerned with much wider contexts than before. (It might be appropriate to remind the reader of similar trends in speech recognition.) It is not possible to take into account all possible coarticulation effects by simply increasing the number of units. At some point the total number might be too high or some units might be based on a very few observations. In this case a normalization of data might be a good solution before the actual unit is chosen. The system will be changed to a rule-based system. However, the rules can be automatically trained from data the same way as in speech recognition [39].

Systems using elements of different lengths depending on the target phoneme and its function are explored by several research groups. In a paper by Olive [38], a new method was described to concatenate "acoustic inventory elements" of different sizes. The system developed at ATR is also based on non-uniform units [44]. These units have been statistically chosen to cover a specific domain.

SPEECH SYNTHESIS IN SPEECH RECOGNITION

One purpose of this paper has been to show how "variability" and "constraints" are relevant aspects to consider in speech synthesis just as in speech recognition. The borders between synthesis and recognition are slowly disappearing. Speaker-independent recognition and speaker-dependent

synthesis have many similar problems to handle. It is, for example, possible to limit the number of input dimensions in a recognition model by using a speaker-dependent synthesis model. Constraints such as vocal tract length, source variation or segment durations can be applied in such a model.

The text-to-speech project at the Royal Institute of Technology has recently focused on modelling different speaker characteristics and speaking styles. Methods in the speech recognition project have been influenced by this work. This has made our speech recognition efforts slightly different from the general trend. The research program, "Nebula" [8], includes prediction models based on speech synthesis. A description of speech on a level closer to articulation, rather than the acoustic base that is used in present-day speech recognition will make generalization of different speakers easier.

Intra-speaker voice source variation can cause severe spectral distortion and contributes to recognition errors in current speaker-independent as well as in speaker-dependent recognition systems. The prosodic information carried by the voice source is important and should not be discarded. This information is lost in many of the current techniques using parameter estimation methods intended to be insensitive to voice source behavior. Since the voice characteristics are changing during an utterance, the speaker adaptation should be part of the recognition process itself. Modelling the source of variation rather than the effect on the speech acoustics potentially makes adaptation more efficient. The production component in the form of a speech synthesis system will ideally make the collection of training data unnecessary. During the last year, special projects studying speaker-independent recognition based on stored phoneme prototypes have been undertaken [6,7]. In these experiments, the references are synthesized during the recognition process itself. The synthetic references can be modified to match the voice of the current speaker. The experiments have shown promising results.

SPEECH SYNTHESIS RESEARCH PRESENTED IN PUBLIC

Speech synthesis research has had a long tradition as a subject for papers and reports. Most work has been presented at meetings such as the ASA or ICASSP. However we can see a tendency towards reduced focus on synthesis papers in these meetings. The European Conference on Speech Technology meetings have shown more variety in synthesis papers. The same is true for the first ICSLP 90 meeting. The successful meeting in Autrans was devoted totally to speech synthesis and showed the breadth of this exciting research area.

One possible reason for lack of publication of detailed work is the reluctance to discuss a particular subject such as "How I improved the synthesis of /s/", "How much f0 movement do I need to turn a statement into a question" or "How I synthesized different degrees of emphasis with a global parameter change". The list just like the research area is endless. Somehow this type of paper is not as acceptable as it used to be. Because of this, some work also stops just before it attains scientific value. At best the result gets hidden in a laboratory system or in some cases in a company product. The message is that the tuning of systems or testing of new solutions must be treated as good research, not as an uninteresting optimization. If we can change this attitude we will get an exciting selection of presentations that will push speech research quality and synthesis quality forward.

CONCLUSION

In this review we have focused on a few exciting research areas which are just beginning to demonstrate their potential. "Speaker variation" and "speaker variability" are key-words in future synthesis research. However we need to go further than just understanding the problem. We need to build new models that can capture the basic parameters along these new dimensions. We need to set up synthesis systems that can incorporate this variation without having to rewrite our software from the beginning. We need to create new synthesizers that can model all the needed control parameters. And we need to structure these parameters in such a way that the users can handle them without too much effort.

To reach these goals we need to make use of methods outside the current speech synthesis domain. Automatic procedures have to be developed to adjust our models to specific speakers and to gather huge amounts of speaker-dependent data. However, we should not get lost in this data. It needs to be structured in new models. Thus the long history of using speech synthesis to evaluate gained knowledge should be continued and expanded [19,50]. Speech synthesis will also be an important research tool in the future.

Acknowledgements

This paper was prepared while the author was a guest researcher at the Spoken Language Systems Group, Laboratory for Computer Science, MIT, USA.

REFERENCES

- [1] Abe, M. (1991) "A segment-based approach to voice conversion," Proc. ICASSP-91.
- [2] Abe, M., Shikano, K. and Kuwabara, H. (1990) "Voice conversion for an interpreting telephone," Proc. Speaker characterization in speech technology, Edinburgh, UK.
- [3] Bailly, G & Laboissi, R. (1990) "Formant trajectories as audible gestures: an alternative for speech synthesis," J. Phonetics 19, No 1.
- [4] Bickley, C., Stevens, K. & Carlson, R. (1991) "Synthesis of manner and voicing continua based on speech production models," J. Acoust. Soc. Am., Vol 89. No4 3SP7.
- [5] Bladon, A., Carlson, R., Granstrom, B., Hunnicutt, S. & Karlsson, I. (1987) "Text-to-speech system for British English, and issues of dialect and style," Proc. European Conference on Speech Technology, Edinburgh, UK.
- [6] Blomberg, M. (1989) "Synthetic phoneme prototypes in a connected-word speech recognition system," Proc. ICASSP-89
- [7] Blomberg, M. (1990) "Adaptation to a speaker's voice in a speech recognition system based on synthetic phoneme references," Proc. ESCA Workshop on Speaker Characterization in Speech Technology, Edinburgh, UK.
- [8] Blomberg, M., Carlson, R., Elenius, K., Granstrom, B. & Hunnicutt, S. (1988) "Word recognition using synthesized reference templates," Proc. Second Symposium on Advanced Man-Machine Interface Through Spoken Language, Hawaii, USA, also in STL-QPSR 2-3/1988.
- [9] Bosch, L. (1990) "Rule extraction for allophone synthesis," Proc. ESCA Workshop on Speech Synthesis, Autrans, France.
- [10] Boves, L. (1990) "Considerations in the design of a multilingual text-to-speech system," J. Phonetics 19, No 1.
- [11] Cahn, J. E. (1990) "The generation of affect in synthesized speech," Journal of the American Voice I/O Society, vol. 8.
- [12] Campbell, N. & Isard, S.D. (1990) "Segment durations in a syllable frame," J. Phonetics 19, No 1.
- [13] Carlson, R. (1991) "Duration models in use," Proc. XIIth ICPHS, Aix en Provence, France.
- [14] Carlson, R., Fant, G., Gobl, C., Granstrom, B., Karlsson, I. & Lin, Q. (1989) "Voice source rules for text-to-speech synthesis," Proc. ICASSP-89.
- [15] Carlson, R., Granstrom, B. & Karlsson, I. (1990) "Experiments with voice modelling in speech synthesis," Proc. ESCA workshop on Speaker Characterization in Speech Technology, Edinburgh, UK.
- [16] Carlson, R., Granstrom, B. & Hunnicutt (1991) "Multilingual text-to-speech development and applications," A.W. Ainsworth (Ed.), Advances in speech, hearing and language processing, JAI Press, London, UK.
- [17] Carpentier, F. & Moulines, E. (1989) "Pitch-synchronous waveform processing techniques for text-to-speech synthesis using diphones," Proc. European Conference on Speech Technology 89.
- [18] Collier, R. (1990) "Multi-language intonation synthesis," J. Phonetics 19, No 1.
- [19] Fant, G. (1990) "Basic research as a support for speech synthesis," J. Phonetics 19, No 1.
- [20] Fant, G., Kruckenberg, A. & Nord, L. (1990) "Prosodic and segmental speaker variations," Proc. ESCA Workshop on Speaker Characterization in Speech Technology, Edinburgh, UK.

- [21] Fant, G., Liljencrants, J. & Lin, Q. (1985) "A four parameter model of glottal flow," *Speech Transmission Laboratory Quarterly and Status Report STL-QPSR* No 4.
- [22] Gopal, H.S., Manzella, J. & Carey, C. (1991) "Factors influencing the spectral representation of front-back vowels in American English," *J. Acoust. Soc. Am.* 4SP10, Vol 89, No4.
- [23] Granstrom, B. & Nord, L. (1991) "Ways of exploring speaker characteristics and speaking styles," *Proc. XIIth ICPhS, Aix en Provence, France.*
- [24] Hakoda K., Nakajima S., Hirokawa T. & Mizuno H. "A new japanese text-to-speech synthesizer based on COC synthesis method," *Proc. ICSLP90, Kobe, Japan.*
- [25] Hertz, S. (1990) "Streams, phones, and transitions: toward a new phonological and phonetic model of formant timing," *J. Phonetics* 19, No 1.
- [26] Holmes W.J. & Pearce D.J.B. "Automatic derivation of segment models for synthesis-by-rule," *Proc. ESCA Workshop on Speech Synthesis, Autrans, France.*
- [27] Huang, C. (1990) "Effects on context, stress, and, speech style on American vowels," *Proc. ICSLP90, Kobe, Japan.*
- [28] Javkin, H. et al. (1989) "A multi-lingual text-to-speech system," *Proc. ICASSP-89.*
- [29] Kaiki, N., Takeda, K. & Sagisaka, Y. (199) "Statistical analysis for segmental duration rules in Japanese speech synthesis," *Proc. Int. Conf. on Spoken Language Processing, Kobe, Japan.*
- [30] Karlsson, I. (1990) "Female voices in speech synthesis," *J. Phonetics* 19, No 1.
- [31] Karlsson, I. (1991) "Dynamic voice quality variations in female speech," *Proc. XIIth, International Congress of Phonetic Sciences, Aix-en-Provence, France.*
- [32] Klatt, D. & Klatt, L. (1990) "Analysis, synthesis, and perception of voice quality variations among female and male talkers," *J. Acoust. Soc. Am.*, vol. 87(2)
- [33] Kohler, K. (1990) "Prosody in speech synthesis: The interplay between basic research and TTS application," *J. Phonetics* 19, No 1.
- [34] Leeuwen, H.C. van & te Lindert, E. (1991) "Speech maker: text-to-speech synthesis based on a multilevel, synchronized data structure," *Proc. ICASSP-91.*
- [35] Lindblom, B. (1990) "Explaining phonetic variation: A sketch of the H&H theory," in *Speech production modelling* Hardcastle and Marchal (eds.) Kluwer Academic Publishers, Netherlands.
- [36] Mouline, E. et al (1990) "A real-time french text-to-speech system generating high quality synthetic speech," *Proc. ICASSP-90.*
- [37] Nakajima, S. & Hamada, H. (1988) "Automatic generation of synthesis units based on context oriented clustering" *Proc. ICASSP-88.*
- [38] Olive, P. (1990) "A new algorithm for a concatenative speech synthesis system using an augmented acoustic inventory of speech sounds," *Proc. ESCA Workshop on Speech Synthesis, Autrans, France.*
- [39] Philips, M., Glass, J. & Zue, V. (1991) "Automatic learning of lexical representations for sub-word unit based speech recognition systems," *Proc. European Conference on Speech Technology* 91.
- [40] Pierrehumbert, J.B. (1987) "The phonetics of English intonation," *Bloomington: IULC.*
- [41] Rahm, M. Kleijn, B. Schroeter, J. (1991) "Acoustic to articulatory parameter mapping using an assembly of neural networks," *Proc. ICASSP-91.*
- [42] Riley, M (1990) "Tree-based modelling for speech synthesis," *Autrans, France.*
- [43] Sagisaka, Y. & Kaiki, N. (1991) "Prosody control for spontaneous speech synthesis," *XIIth, International Congress of Phonetic Sciences, Aix-en-Provence, France.*
- [44] Sagisaka, Y. (1988) "Speech synthesis by rule using an optimal selection of non-uniform synthesis units," *Proc. ICASSP -88.*
- [45] Santen, J. van & Olive, J. (1990) "The analysis of segmental effect on segmental duration," *Computer Speech and Language*. No 4.
- [46] Scherer, K. (1989) "Vocal correlates of emotion," in (eds. Wagner, H. & Manstead, T.), *Handbook of Psychophysiology: Emotion and Social Behavior*, Chichester: Wiley, UK.
- [47] Son, R.J.J.H. van & Pols, L. (1989) "Comparing formant movements in fast and normal rate speech," *Proc. European Conference on Speech Technology* 89.
- [48] Sorin, C., Larreur, D. & Llorca, R. (1987) "A rhythm-based prosodic parser for text-to-speech systems in French," *Proc. XIth ICPhS, Tallin, Estonia.*
- [49] Stevens, K. & Bickley, C. (1990) "Constraints among parameter simplify control of Klatt formant synthesizer," *J. Phonetics* 19, No 1.
- [50] Stevens, K. (1991) "The contribution of speech synthesis to phonetics: Dennis Klatt's legacy," *Proc. XIIth, International Congress of Phonetic Sciences, Aix-en-Provence, France.*
- [51] Talkin D. & Rowley J. "Pitch-synchronous analysis and synthesis for TTS systems," *Proc. ESCA Workshop on Speech Synthesis, Autrans, France.*
- [52] Williams, C. E. & Stevens, K. N. (1972) "Emotions and speech: some acoustical correlates," *J. Acoust. Soc. Am.* Vol. 52.