

June 22, 1965

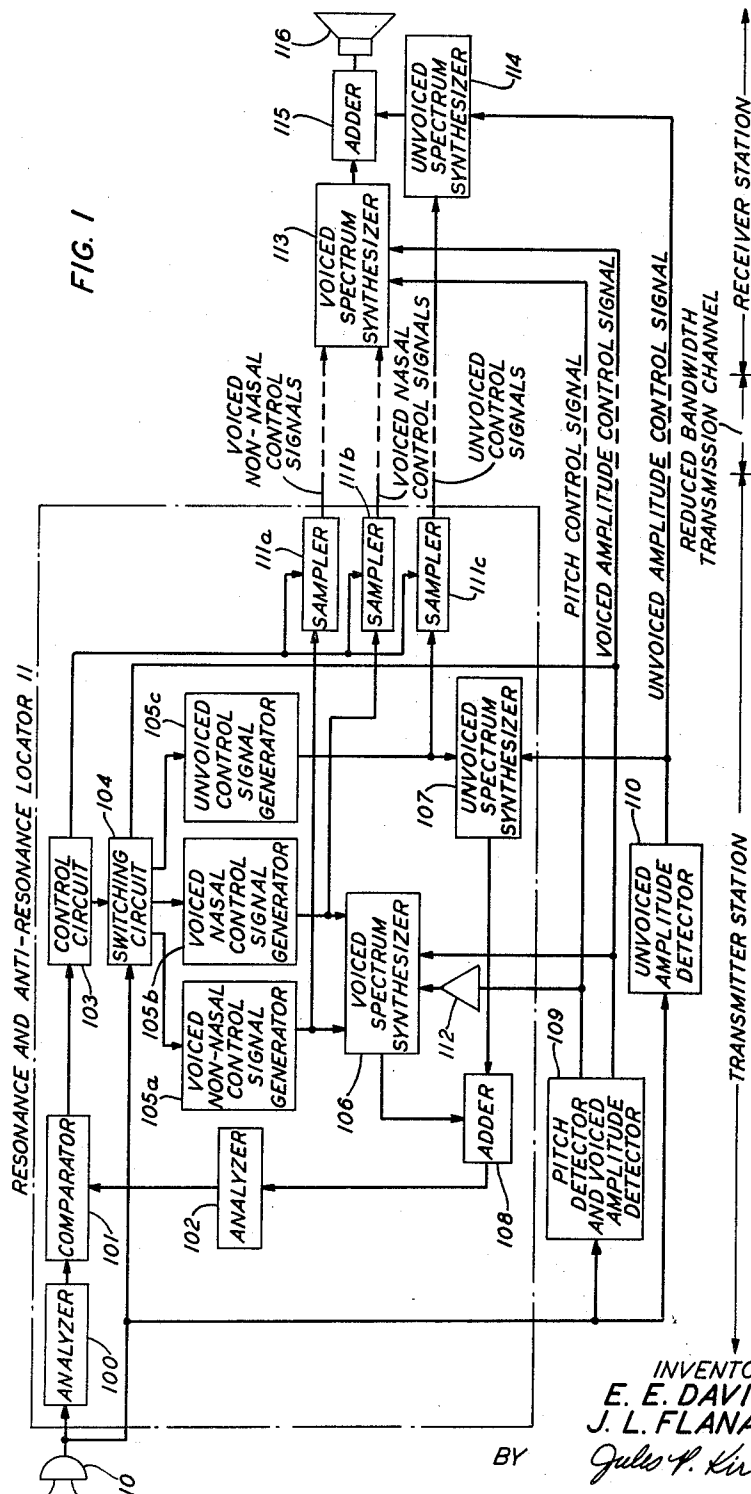
E. E. DAVID, JR., ETAL

3,190,963

TRANSMISSION AND SYNTHESIS OF SPEECH

Filed Nov. 6, 1962

7 Sheets-Sheet 1



June 22, 1965

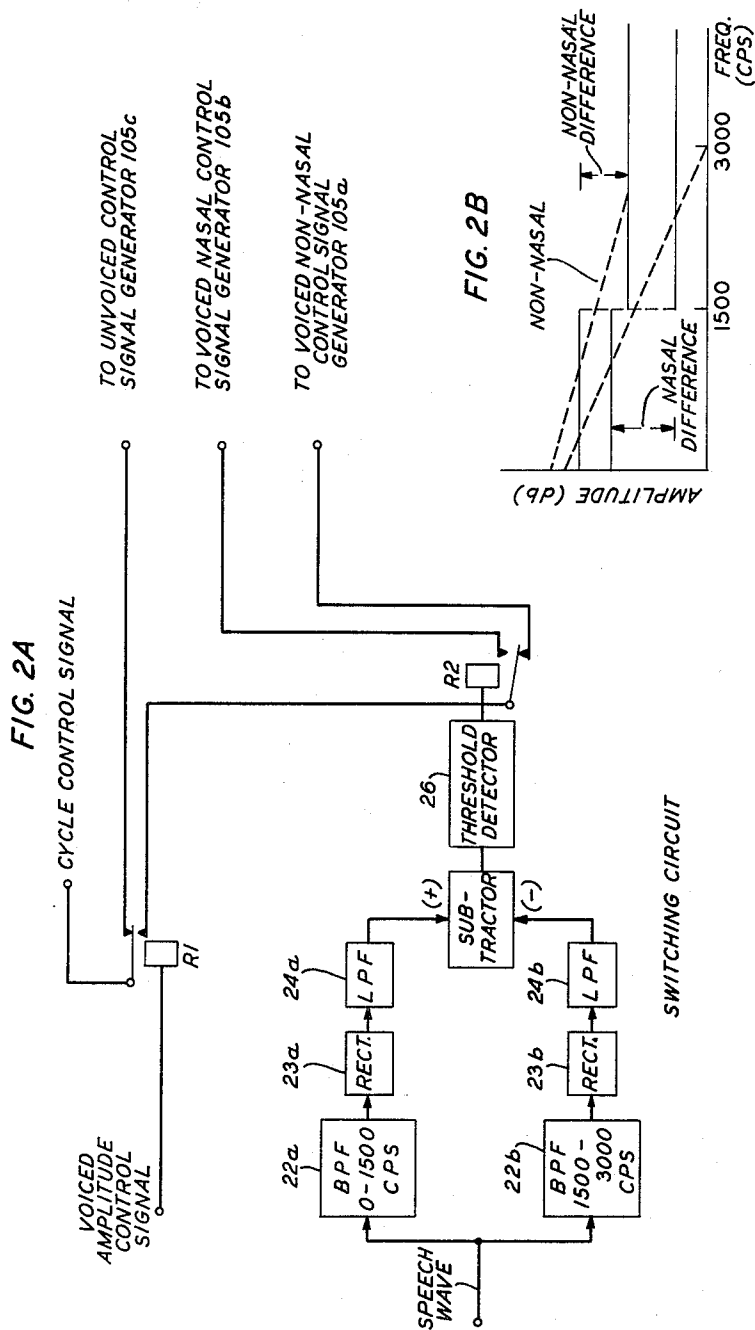
E. E. DAVID, JR., ETAL

3,190,963

TRANSMISSION AND SYNTHESIS OF SPEECH

Filed Nov. 6, 1962

7 Sheets-Sheet 2



June 22, 1965

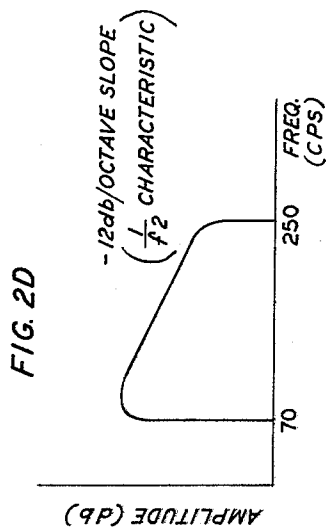
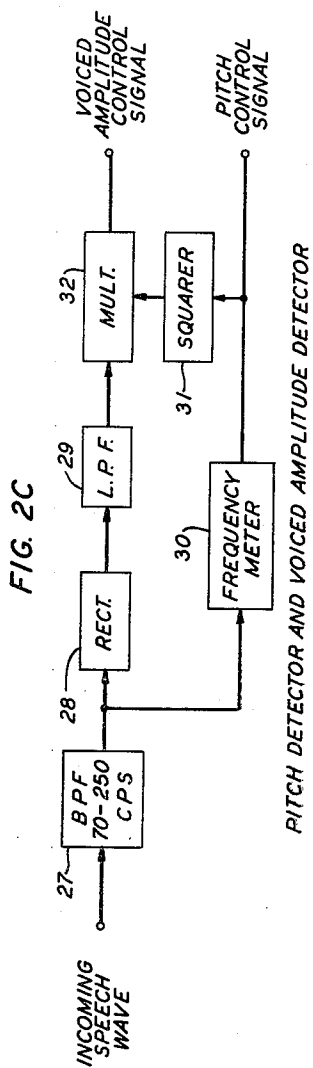
E. E. DAVID, JR., ETAL

3,190,963

TRANSMISSION AND SYNTHESIS OF SPEECH

Filed Nov. 6, 1962

7 Sheets-Sheet 3



June 22, 1965

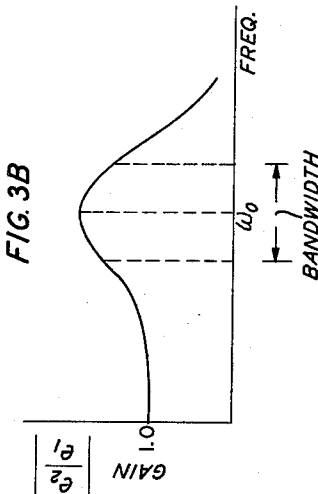
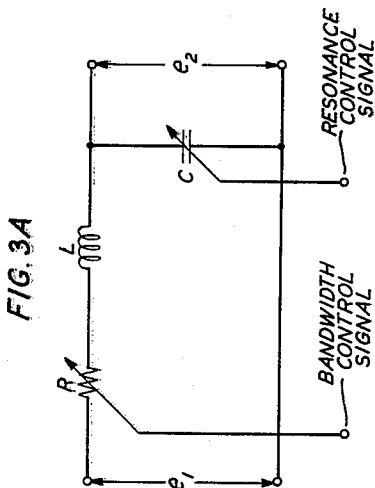
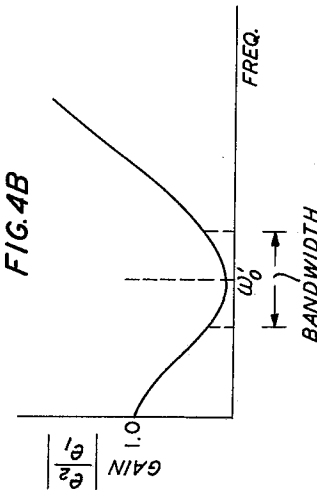
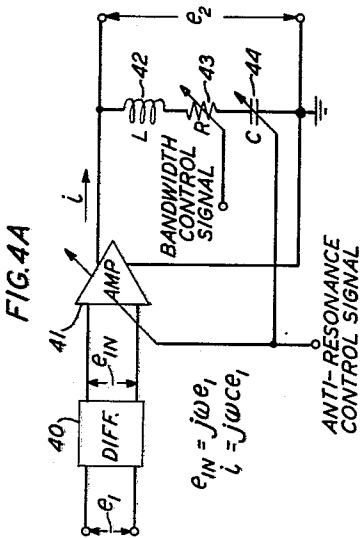
E. E. DAVID, JR., ETAL

3,190,963

TRANSMISSION AND SYNTHESIS OF SPEECH

Filed Nov. 6, 1962

7 Sheets-Sheet 4



June 22, 1965

E. E. DAVID, JR., ETAL

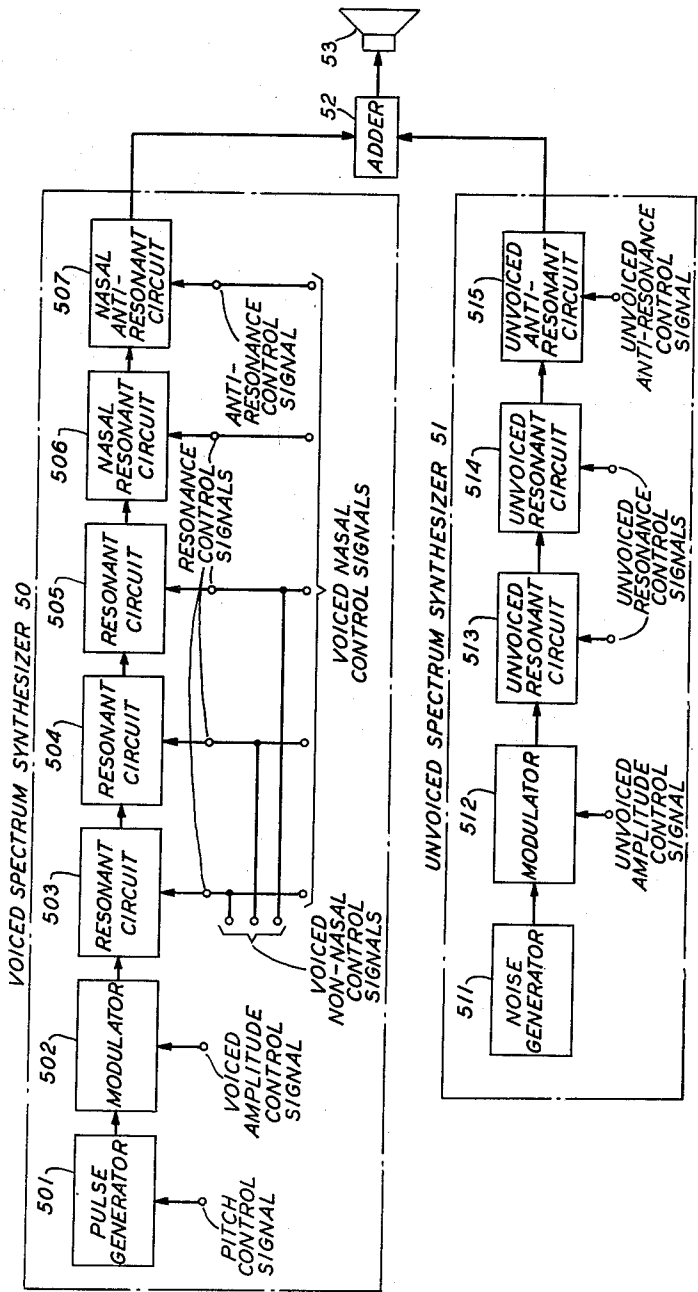
3,190,963

TRANSMISSION AND SYNTHESIS OF SPEECH

Filed Nov. 6, 1962

7 Sheets-Sheet 5

FIG. 5



June 22, 1965

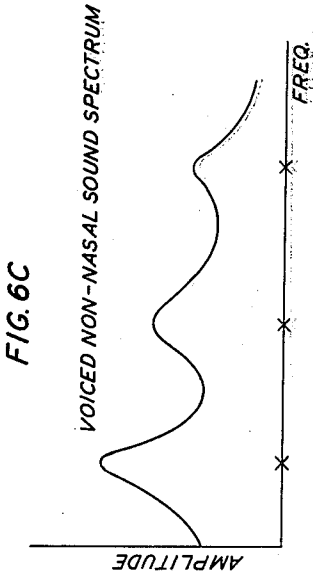
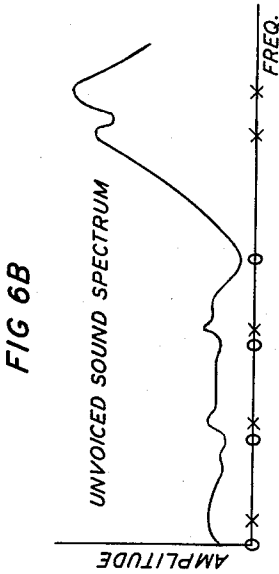
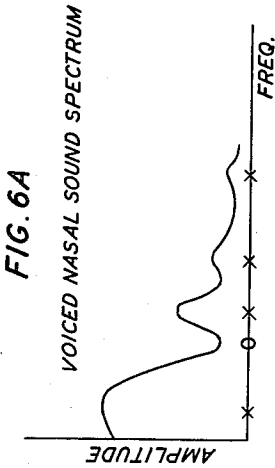
E. E. DAVID, JR., ETAL

3,190,963

TRANSMISSION AND SYNTHESIS OF SPEECH

Filed Nov. 6, 1962

7 Sheets-Sheet 6



LEGEND: X DENOTES RESONANCE
O DENOTES ANTI-RESONANCE

June 22, 1965

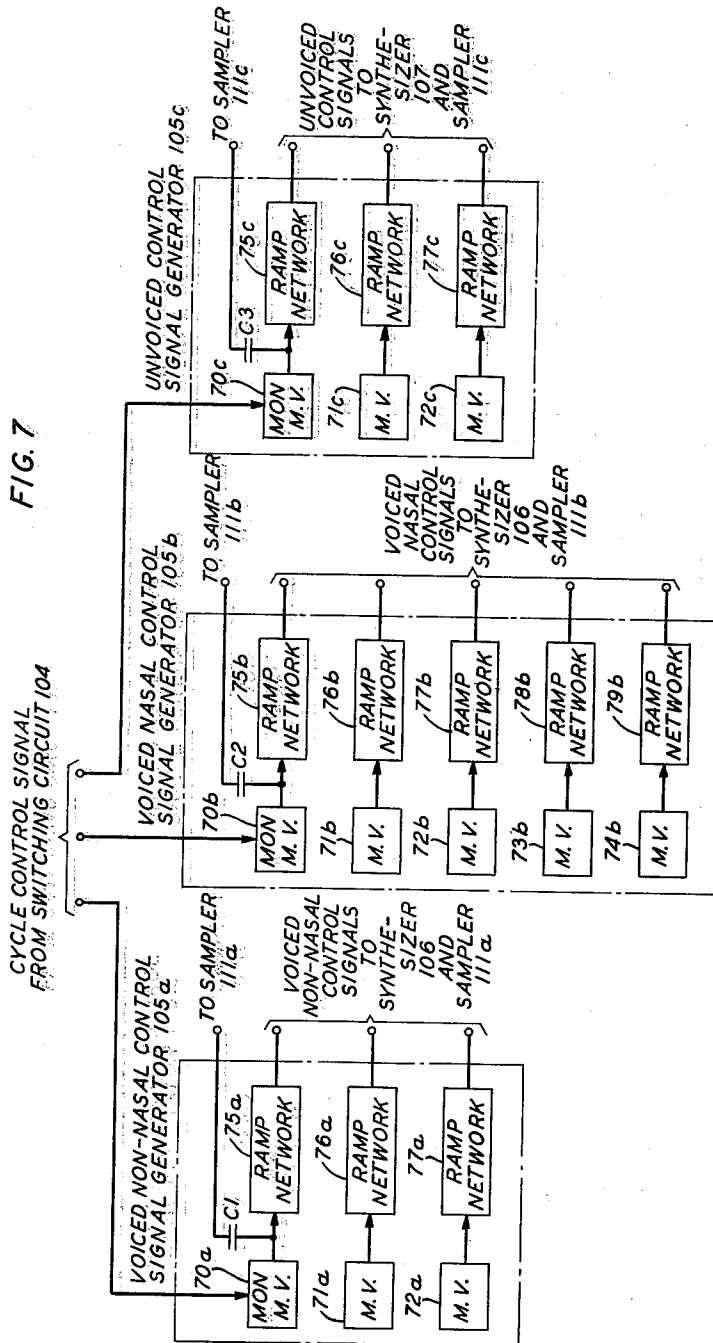
E. E. DAVID, JR., ETAL

3,190,963

TRANSMISSION AND SYNTHESIS OF SPEECH

Filed Nov. 6, 1962

7 Sheets-Sheet 7



1

3,190,963

TRANSMISSION AND SYNTHESIS OF SPEECH

Edward E. David, Jr., Berkeley Heights, and James L. Flanagan, Warren Township, Somerset County, N.J., assignors to Bell Telephone Laboratories, Incorporated, New York, N.Y., a corporation of New York

Filed Nov. 6, 1962, Ser. No. 235,703

9 Claims. (Cl. 179—15.55)

This invention relates to the transmission and synthesis of speech, and in particular to the transmission and synthesis of speech in bandwidth compression systems.

In order to make more economic use of the frequency bandwidth of speech transmission channels, a number of bandwidth compression arrangements have been devised for transmitting the information content of a speech wave over a channel whose bandwidth is substantially narrower than that required for transmission of the speech wave itself. Bandwidth compression systems typically include at a transmitting terminal an analyzer for deriving from an incoming speech wave a group of narrow bandwidth control signals representative of selected information bearing characteristics of the speech wave, and at a receiving terminal, a synthesizer for reconstructing from the control signals a replica of the original speech wave.

One well-known bandwidth compression system is the so-called "resonance vocoder," a specific form of which is described in E. S. Weibel Patent 2,817,707, issued December 24, 1957. In a resonance vocoder, the information bearing characteristics which are represented by the control signals and which are reconstructed at the receiving terminal are the locations in the speech amplitude spectrum of selected resonant or formant frequencies. These resonances or formants correspond to the principal vocal tract resonances, that is, they correspond to frequency regions of relatively effective transmission through a talker's vocal tract.

It has been determined, however, that although resonances describe a large and important class of speech sounds, the so-called voiced nonnasal sounds, there are at least two other large and important classes of speech sounds which are not adequately characterized by resonant frequencies alone, namely, voiced nasal sounds and unvoiced sounds. In particular, it has been recognized that an adequate description of the two latter classes of sounds requires the location of at least one principal antiresonant frequency in addition to the locations of selected resonant frequencies, where antiresonant frequencies correspond to regions of relatively ineffective transmission through a talker's vocal tract.

Recognition of the important role played by antiresonant frequencies in the accurate specification of voiced nasal sounds and unvoiced sounds has been turned to account in the bandwidth compression system of the present invention. At the transmitting terminal of this invention, an incoming speech wave is analyzed to determine the locations of selected resonant and antiresonant frequencies. These frequencies are represented by reduced bandwidth control signals that are transmitted to a receiving terminal where they are utilized to reconstruct a natural sounding replica of the original speech wave.

The location of an antiresonance is evidenced by a major valley or dip in the envelope of the amplitude spectrum of a particular sound, and in the present invention an antiresonance is determined by locating the corresponding

2

valley in the spectral envelope. On the other hand, it is well known that locations of resonances correspond to major peaks in the spectral envelope of a sound, and resonances are determined in this invention by locating these peaks in the spectral envelope.

Peaks and valleys in the spectral envelope of an incoming speech wave are located at the analyzer of the present invention by comparing one of three locally generated, artificial spectra with the speech spectrum. The three artificial spectra correspond to the three classes of speech sounds, voiced nonnasal sounds, voiced nasal sounds, and unvoiced sounds, and selection of the appropriate artificial spectrum depends upon which of the three classes of speech sounds appears in the speech wave at a given instant. Each artificial spectrum is constructed from a set of resonance or resonance and antiresonance signals, and each set of resonance or resonance and antiresonance signals is varied cyclically over a wide range of values corresponding to a wide range of resonance or resonance and antiresonance locations. In this fashion, the comparison of the speech spectrum with an artificial spectrum during each cycle of variation produces a best matching version of the artificial spectrum whose resonances or resonances and antiresonances correspond closely to the true speech resonances or resonances and antiresonances. The best matching resonances or resonances and antiresonances for each of the three classes of speech sounds are represented by one of three groups of narrow band control signals which are successively transmitted to a receiving terminal, and at the receiving terminal there signals a corresponding succession of accurate replicas of the spectra of the succession of different sounds in the original speech wave. By constructing each replica spectrum at the receiving terminal to correspond to one of three above-mentioned major speech sound classes, a high degree of realism for a large number of sounds is attained in the speech reproduced at the receiving terminal.

The invention will be fully understood from the following descriptions of illustrative embodiments thereof, taken in connection with the appended drawings, in which:

FIG. 1 is a block diagram showing the complete bandwidth compression system of this invention;

FIGS. 2A and 2C are block diagrams illustrating in detail certain components of the system shown in FIG. 1;

FIGS. 2B and 2D are graphs illustrating particular features of the circuits shown in FIGS. 2A and 2C;

FIGS. 3A and 4A are schematic diagrams of circuits for synthesizing resonances and antiresonances, respectively;

FIGS. 3B and 4B are graphs of assistance in explaining the operation of the circuits shown in FIGS. 3A and 4A, respectively;

FIG. 5 is a block diagram showing apparatus for synthesizing speech sounds in accordance with the principles of this invention;

FIGS. 6A, 6B, and 6C illustrate spectral envelopes of typical voiced nasal, unvoiced, and voiced nonnasal sounds, respectively; and

FIG. 7 is a block schematic diagram showing in detail certain components of the system illustrated in FIG. 1.

Theoretical considerations

As explained in detail by J. L. Flanagan in "A Resonance-Vocoder and Baseband Complement: A Hybrid System for Speech Transmission," 1959 Institute of Radio

Engineers Wescon Convention Record, part 7, pages 5-16, speech sounds may be divided into three classes, voiced nonnasal, voiced nasal, and unvoiced (nonnasal), each of which is characterized by a so-called vocal transmission function. In the case of voiced nonnasal sounds, the variable features of the corresponding vocal transmission function may be specified solely in terms of the first few poles or resonances of the human vocal tract, whereas in the case of voiced nasal sounds and unvoiced sounds, the variable features of the corresponding vocal transmission functions must be specified in terms of zeros or antiresonances as well as poles or resonances of the human vocal tract. As explained in G. Fant, *Acoustic Theory of Speech Production* (1960), the production of human speech may be characterized in several different ways. Thus the production of human speech may be described in terms of the human vocal mechanism which includes the glottis, the vocal tract, the nasal cavities, and other articulatory features. The vocal tract, for example, resembles a tube of non-uniform cross section which is characterized by a number of resonances or frequency regions of relatively effective transmission through the vocal tract. Another description of speech production is in terms of the properties of the speech sound wave, which has an amplitude spectrum characterized by a number of prominent peaks or formants that correspond to resonances of the vocal tract, as pointed out on page 20 of the Fant reference. Finally, by analogy with the terminology of acoustic and electrical engineering, the production of speech may be described as the response of the vocal tract filter systems to one or more sound sources, in which case the vocal tract is characterized by a transmission function that may be expressed as a rational function of the complex frequency variable, s , as shown on page 42 of the Fant text. This rational function expression has both zeros and poles, with the zeros corresponding to antiresonances of the vocal tract and the poles corresponding to resonances of the vocal tract, as pointed out on pages 48 through 62 of the Fant reference.

Resonances and antiresonances of the various classes of sounds are manifested by maxima and minima, respectively, in the envelopes of the amplitude spectra of these sounds, as pointed out on page 24 of the Fant reference cited above, and as shown and described by J. L. Flanagan in "Note on the Design of 'Terminal-Analog' Speech Synthesizers," vol. 29, *Journal of the Acoustical Society of America*, pages 306, 310 (1957). FIGS. 6A, 6B, and 6C illustrate the envelopes of the amplitude spectra of typical voiced nasal, unvoiced, and voiced nonnasal sounds, respectively, in which the symbol "x" on the horizontal axis indicates the approximate frequency location of a spectral maximum or resonance, and the symbol "O" on the horizontal axis indicates the approximate frequency location of a spectral minimum or antiresonance.

It is observed in FIG. 6C that there are three principal maxima or resonances in the spectrum of a typical voiced nonnasal sound, while in FIG. 6A there are four principal maxima or resonances and a single principal minimum or antiresonance in the spectrum of a typical voiced nasal sound. The three resonances in FIG. 6C are characteristic of the role of the principal vocal tract resonances in the production of voiced nonnasal sounds, while in FIG. 6A the first three resonances are characteristic of the principal nasal tract resonances. The additional resonance and antiresonance in FIG. 6A which do not appear in FIG. 6C are characteristic of the role of the vocal tract in the production of voiced nasal sounds, that is, in the production of voiced nasal sounds, the nasal tract is the main branch for the flow of air from the lungs, and the vocal tract acts as a side branch. For voiced nonnasal sounds of the type illustrated by FIG. 6C, the vocal tract acts as the main branch.

In the unvoiced sound spectrum shown in FIG. 6B it is noted that there are four resonances and four anti-

resonances, but that there are three pairs of resonances and antiresonances in which the resonance and antiresonance lie so close together that they essentially nullify each other, as evidenced by the small maxima and minima corresponding to these resonance-antiresonance pairs in the spectrum. Hence unvoiced sounds may be specified in terms of two resonances and a single antiresonance.

In the following description of the apparatus of this invention, the three classes of sounds described above will be characterized in terms of their principal resonances and antiresonances as shown in FIGS. 6A, 6B, and 6C. However, it is to be understood that these sounds may be specified by additional resonances and antiresonances, if desired, with appropriate modification of the necessary components of this invention.

Complete system

Referring first to FIG. 1, this drawing illustrates a complete bandwidth compression system in which both resonances and antiresonances are utilized as the information bearing characteristics of speech to conserve transmission channel bandwidth. At the transmitter station of this system, an incoming speech wave from source 10, for example, a conventional microphone, is applied simultaneously to resonance and antiresonance locator 11, pitch detector and voiced amplitude detector 109, and unvoiced amplitude detector 110. As described in detail below, locator 11 produces at any given time one of three groups of narrow band control signals, one representative of selected resonant frequencies of voiced nonnasal sounds, one representative of selected resonant and antiresonant frequencies of voiced nasal sounds, and the other representative of selected resonant and antiresonant frequencies of unvoiced sounds. Thus the succession of these three different classes of sounds in the incoming speech wave causes locator 11 to produce a corresponding succession of groups of control signals each representative of the corresponding sound. Detector 109, which is illustrated in detail in FIG. 2C, derives from the speech wave a pitch control signal indicative of the frequency of the fundamental component of voiced portions of the speech wave, and a voiced amplitude control signal representative of the amplitude of the fundamental component of voiced portions of the speech wave. Detector 110, which may comprise a conventional rectifier followed by a low pass filter, develops an unvoiced amplitude control signal indicative of the energy of unvoiced portions of the speech wave.

The control signals produced at the transmitter station specify the characteristics of the three classes of speech sounds referred to above with a high degree of accuracy, and the total frequency bandwidth of these control signals is substantially smaller than that of the speech wave from source 10. Hence these control signals may be transmitted to a distant receiver station over a reduced bandwidth transmission channel, a suitable channel being indicated in FIG. 1 by broken lines.

At the receiver station, the voiced nasal and nonnasal control signals, the pitch control signal, and the voiced amplitude control signal are applied to voiced spectrum synthesizer 113 which reconstructs replicas of the spectra of voiced nasal and nonnasal sounds. The unvoiced resonance and antiresonance control signals and the unvoiced amplitude control signal are delivered to unvoiced spectrum synthesizer 114, which reconstructs replicas of the spectra of unvoiced sounds. The structure of synthesizers 113 and 114 is illustrated in FIGS. 3A, 4A, and 5, and the reconstructed spectra developed by these synthesizers are combined in adder 115, which may be any one of a number of well-known adding circuits. The combined spectra are converted into audible speech sounds by reproducer 116, for example, a conventional loudspeaker.

Resonance and antiresonance locator 11 is an improvement upon the formant locator apparatus shown in Patent

3,127,477 issued March 31, 1964, to E. E. David, Jr., M. V. Mathews, and J. E. Miller, Serial No. 205,663, filed June 27, 1962, in that locator 11 determines both resonances and antiresonances in the spectrum of an incoming speech wave, whereas the formant locator apparatus shown in the David et al. patent determines only resonances or formants. The process of determining resonances and antiresonances embodied by locator 11 comprises a rapid succession of repetitive cycles initiated by a repetitive cycle control signal generated by control circuit 103 and delivered through switching circuit 104 to one of the three signal generators 105a, 105b, or 105c. Control circuit 103 may be similar in design to control circuit 23 of the above-mentioned E. E. David, Jr., et al. patent but switching circuit 104 is shown in detail in FIG. 2A of the present application.

Switching circuit 104 is supplied with the incoming speech wave from source 10, and after determining which one of the three classes of sound is present in the speech wave, circuit 104 passes the cycle control signal to the appropriate signal generator. In response to the cycle control signal passed by switching circuit 104, the appropriate signal generator repetitively produces in each cycle a set of signals from which an artificial spectrum may be constructed. Voiced nonnasal control signal generator 105a produces a set of signals representing selected resonant frequencies of voiced nonnasal sounds, voiced nasal control signal generator 105b produces a set of signals representing selected resonant and antiresonant frequencies of voiced nasal sounds, and unvoiced control signal generator 105c produces a set of control signals representing selected resonant and antiresonant frequencies of unvoiced sounds. Each of the signals produced by generators 105a through 105c varies continuously over a range of values corresponding to the usual frequency range of a particular resonance or antiresonance. Further, the combination of values represented by each set of control signals is made to vary in a manner similar to that shown in the aforementioned E. E. David, Jr., et al. patent so that during each cycle of operation of locator 11, each set of control signals collectively represents all possible combinations of resonance or resonance and antiresonance locations.

The sets of control signals from generators 105a and 105b are applied to voiced spectrum synthesizer 106, together with the pitch and voiced amplitude control signals from detector 109. Synthesizer 106 is similar in design to synthesizer 113 at the receiver station, illustrated in detail and labeled synthesizer 50 in FIG. 5, and constructs from the incoming control signals an artificial signal having an amplitude spectrum whose resonances or resonances and antiresonances vary continuously in location in synchrony with the continuous variations in value of one or the other set of applied control signals. Similarly, the set of control signals from generator 105c is applied to unvoiced spectrum synthesizer 107, together with the unvoiced amplitude control signal from detector 110. Synthesizer 107 may be identical in structure with synthesizer 114 at the receiver station, illustrated in detail and labeled synthesizer 51 in FIG. 5, and constructs from the incoming control signals an artificial signal having an amplitude spectrum whose resonances and antiresonances vary continuously in location in synchrony with the continuous variations in value of the set of applied control signals.

In order to make the spectra of the artificial signals from synthesizers 106 and 107 resemble the spectra of the various sounds contained in the speech wave in every important respect save resonance and antiresonance locations, it is necessary that the frequency components of the artificial spectra be scaled up in frequency, as explained in the previously mentioned E. E. David, Jr., et al. patent. This is accomplished for the artificial signals produced by synthesizer 106 by passing the pitch control signal from detector 109 through a conventional voltage

amplifier 112 having an appropriate gain constant. For the artificial signals produced by synthesizer 107 this is achieved by providing synthesizer 107 with a noise source that generates frequency components with sufficiently high frequencies; that is, as shown in FIG. 5, noise generator 511 of synthesizer 51 may be a conventional wide band gas diode source.

The artificial signals produced by synthesizers 106 and 107 are combined in adder 108, which may be of any well-known sort, to form at the output terminal of adder 108 a succession of artificial signals having spectra corresponding to the succession of different classes of sounds appearing in the incoming speech wave. From adder 108 the succession of artificial signals is passed to analyzer 102, which separates each successive artificial signal into the individual frequency components that constitute its amplitude spectrum. Analyzer 100 similarly separates the incoming speech wave from source 10 into the individual frequency components that constitute the speech spectrum, and the two sets of frequency components from analyzers 100 and 102 are sent to comparator 101. Analyzers 100 and 102 and comparator 101 may each be similar in construction with analyzers 21 and 27 and comparator 22, respectively of the above-mentioned E. E. David, Jr., et al. patent.

Comparator 101 operates upon the two sets of frequency components from analyzers 100 and 102 to derive an error signal whose magnitude is representative of the difference between each artificial spectrum and the speech spectrum during each cycle. The difference between the two spectra, as represented by the error signal, arises from the difference between the resonance or resonance and antiresonance locations of the speech spectrum and the respective resonance or resonance and antiresonance locations of the artificial spectrum, because in every other important respect, that is, in pitch and in amplitude as derived by detectors 109 and 110, the speech spectrum and the artificial spectrum are identical.

Since the resonances or resonances and antiresonances of each artificial spectrum are made to occur at substantially all possible resonance or resonance and antiresonance locations during each cycle, theoretically there will be at least one point in each cycle when the resonance or resonance and antiresonance locations of the artificial spectrum will be identical with the resonance or resonance and antiresonance locations of the speech spectrum. Accordingly, the magnitude of the error signal developed by comparator 101 at such a point in each cycle would be zero. However, because of noise and other imperfections, the artificial spectrum will never be exactly identical with the speech spectrum, but there is at least one point in each cycle at which the difference between the two spectra is at a minimum, this minimum difference being indicated by a corresponding minimum in the magnitude of the error signal produced by comparator 101.

To determine the resonance or resonance and antiresonance locations of the artificial spectrum at the point in each cycle at which the difference between the artificial spectrum and the speech spectrum is a minimum, the error signal from comparator 101 is passed to control circuit 103, where the error signal is continuously examined to determine its minimum magnitude in each cycle. At the instant in each cycle that circuit 103 detects a minimum magnitude, it delivers a sample control signal to samplers 111a through 111c, which may be identical with sampler 28 of the above-mentioned E. E. David, Jr., et al. patent, and in response, the sampler that is receiving control signals from its corresponding signal generator 105a, 105b, or 105c samples and stores the values of the set of resonance or resonance and antiresonance signals which is being applied from the signal generator. When the cycle ends, these sampled and store values are converted into control signals that indicate with a high degree of accuracy the locations of the resonances or resonances and anti-

resonances of the spectrum of the class of sound present at a given instant in the incoming speech wave.

As described in the E. E. David, Jr., et al. patent referred to above, the difference between an artificial spectrum and the speech spectrum, as represented by the error signal, may be measured in several ways. It is to be understood that any one of these measures of error may be utilized in this invention.

Pitch detector and voiced amplitude detector

Turning now to FIG. 2C, this drawing illustrates the structure of detector 109 employed at the transmitter station of this invention, as shown in FIG. 1. The incoming speech wave from source 10 is passed through bandpass filter 27 to obtain the speech frequency components lying in the frequency range between 70 and 250 cycles per second. Within this frequency range, the bandpass characteristic of filter 27 is designed in conventional fashion to decrease with frequency at the rate of 12 decibels per octave, as shown in FIG. 2D; that is, the amplitude of each component in this frequency range is reduced by a factor of

$$\frac{1}{f^2}$$

where f denotes frequency. Since the speech frequency components lying within this pass band have a natural tendency to decrease in amplitude with increasing frequency at a rate of 12 decibels per octave, the amplitude of the first or fundamental frequency component of the speech wave is further enhanced by passage through filter 27.

Because the chief constituent of the output signal of filter 27 is the fundamental frequency component, a pitch control signal representative of the fundamental speech frequency may be obtained from this output signal. For example, a conventional frequency meter 30 may be employed to determine the reciprocals of the intervals between successive positive-going axis crossings of the output signal of filter 27. The pitch control signal developed by frequency meter 30 is then utilized in the manner shown in FIG. 1.

A signal representative of the magnitude of the fundamental speech frequency component voiced sounds may also be derived from the output signal of filter 27. This is accomplished by rectifying the output signal of filter 27 and averaging the rectified signal over an interval comparable to the period of the fundamental speech component, for example, over an interval on the order of 10 to 20 milliseconds. A conventional rectifier 28 and low pass filter 29 may be utilized to perform the rectifying and averaging operations. However, the magnitude of the rectified and averaged signal developed at the output terminal of filter 29 is not indicative of the amplitude of the fundamental frequency component because the characteristic of filter 27 reduced the magnitude of the fundamental component by a factor

$$\frac{1}{f^2}$$

To cause the magnitude of the rectified and averaged output signal of filter 29 to represent the magnitude of the fundamental speech component, it is therefore necessary to increase the magnitude of the rectified and averaged signal by a factor f^2 , for example, by squaring the magnitude of the pitch control signal and multiplying the rectified and averaged signal by the squared pitch control signal. These operations may be realized in the fashion indicated in FIG. 2C by squaring circuit 31 and multiplier 32, both of which may be of well-known construction. The output signal of multiplier 32 is then suitable for use as a voiced amplitude control signal, as illustrated in FIG. 1.

Switching circuit

Referring now to FIG. 2A, this drawing shows the structure of switching circuit 104 employed in locator 11 at the transmitter terminal of the system of FIG. 1. The cycle control signal from control circuit 103 is directed by relays R1 and R2 to one of the three signal generators 105a, 105b, or 105c according to the following logic. Relay R1, which is provided with a control terminal, an input terminal, and two output terminals, may be any one of a number of well-known devices for directing an incoming signal from its input terminal to one or the other of its output terminals in response to whether or not a control signal of sufficient magnitude is applied to its control terminal. In FIG. 2A, the voiced amplitude control signal is applied to the control terminal of relay R1, the cycle control signal is applied to the input terminal of relay R1, and the output terminals of relay R1 are connected to unvoiced signal generator 105c and to the input terminal of relay R2. Relay R1 is energized by the voiced amplitude control signal from detector 109 so that the cycle control signal is passed to the input terminal of relay R2 whenever a voiced sound, either nasal or nonnasal, is present in the incoming speech wave. When an unvoiced sound is present in the incoming speech wave, the magnitude of the voiced amplitude control signal falls below a predetermined level necessary to energize relay R1, and in its deenergized state, relay R1 passes the cycle control signal to unvoiced control signal generator 105c.

To understand the operation of circuit 104 in distinguishing between voiced nasal and voiced nonnasal sounds it is convenient at this point to refer to the graph in FIG. 2B. The upper dashed line in the graph of FIG. 2B illustrates the relatively small decrease of amplitude as a function of frequency for the spectrum of a typical nonnasal voiced sound, while the lower dashed line illustrates the relatively large decrease of amplitude as a function of frequency for the spectrum of a typical nasal voiced sound. It is therefore observed that in the frequency range from 0 to 3,000 cycles per second the difference in average amplitude between the components in the range from 0 to 1,500 cycles per second and the components in the range from 1,500 to 3,000 cycles per second is greater for nasal voiced sounds than for nonnasal voiced sounds. This characteristic difference in average amplitude is turned to account in the switching circuit apparatus shown in FIG. 2A to direct the cycle control signal to one of the two signal generators 105a or 105b.

The incoming speech wave from source 10 is applied simultaneously to two parallel paths, each path containing a bandpass filter, a rectifier, and a low pass filter connected in series. The bandpass filter 22a in the upper path is proportioned to pass speech components lying in the frequency range from 0 to 1,500 cycles per second, while bandpass filter 22b in the lower path is proportioned to pass speech components within the frequency range from 1,500 to 3,000 cycles per second. The output terminal of the upper path is connected to the minuend terminal of a conventional subtractor device 25, and the output terminal of the lower subpath is connected to the subtrahend terminal of subtractor 25. The magnitude of the difference signal developed at the output terminal of subtractor 25 is indicative of whether a nasal or a nonnasal voiced sound is present in the incoming speech wave. The difference signal from subtractor 25 is passed to a threshold detector 26, which may be of any well-known design, and if the magnitude of the difference signal exceeds the predetermined threshold of detector 26, thereby indicating the presence of a nasal voiced sound, an output signal from detector 26 is applied to the control terminal of relay R2. When relay R2 is energized by a control signal from detector 26, the incoming cycle control signal is passed to signal generator 105b. If the magnitude of the difference signal does not exceed the predetermined threshold of detector 26, thereby indicat-

ing the presence of a nonnasal voiced sound, relay R2 remains deenergized and the cycle control signal is delivered to signal generator 105a.

Signal generators 105a, 105b, 105c

The structure of signal generators 105a through 105c is based upon the structure of formant generator 24 shown in the above-mentioned E. E. David, Jr., et al. patent, modified, however, to take into account antiresonances where necessary. Thus in FIG. 7 signal generator 105a may be identical in construction with formant generator 24 illustrated in FIG. 2B of the E. E. David, Jr., et al. patent, that is, signal generator 105a comprises a monostable multivibrator 70a and free running multivibrators 71a and 72a, each of the two latter multivibrators generating a succession of positive-going pulses at rates on the order of 200 and 2,000 cycles per second, respectively. Monostable multivibrator 70a is triggered to its unstable state by the cycle control signal delivered from control circuit 103 through switching circuit 104, and the positive-going pulse generated by multivibrator 70a at the beginning of each cycle is applied through capacitor C1 to sampler 111.

Each multivibrator 70a, 71a, 72a is connected to a ramp network 75a, 76a, 77a, respectively, which converts each positive-going output pulse into a ramp-shaped signal that increases continuously over a range of values corresponding to the range of frequencies within which a particular speech resonance is defined to occur. The ramp networks of signal generator 105a may also be identical in structure with the ramp networks described in the E. E. David, Jr., et al. patent, and during a single cycle the output signals of ramp networks 75a through 77a represent substantially all possible locations of the three principal resonances of voiced nonnasal sounds.

Signal generators 105b and 105c operate upon the same principles as signal generator 105a, but the set of signals generated by each of the two former circuits represents substantially all possible locations of the principal resonances and the principal antiresonance of voiced nasal sounds and unvoiced sounds, respectively. Thus signal generator 105b is provided with one monostable multivibrator 70b, four free running multivibrators 71b through 74b, and five corresponding ramp networks 75b through 79b, to produce a set of five voiced nasal control signals. Suitable frequencies of oscillation for multivibrators 71b through 74b are on the order of 200, 2,000, 20,000, and 200,000 cycles per second, respectively. Four of the five control signals represent the ranges of frequencies within which four of the principal resonances of voiced nasal sounds are defined to occur, and the fifth control signal represents the range of frequencies within which one of the principal antiresonances of voiced nasal sounds is defined to occur, for example, as shown in FIG. 6A.

Similarly, signal generator 105c is provided with one monostable multivibrator 70c, two free running multivibrators 71c and 72c, and three corresponding ramp networks 75c through 77c, to produce a set of three unvoiced control signals. The frequencies of oscillation for multivibrators 71c and 72c may be the same as those of multivibrators 71a and 72a in generator 105a. Two of the three control signals represent the ranges of frequencies within which two of the principal resonances of unvoiced sounds are defined to occur, and the third control signal represents the range of frequencies within which one of the principal antiresonances of unvoiced sounds is defined to occur, for example, as shown in FIG. 6B.

Spectrum synthesizers

Synthesizers suitable for reconstructing the spectra of voiced nasal, voiced nonnasal, and unvoiced sounds at both the transmitter and receiver terminals of this invention are illustrated in detail in FIG. 5. Voiced spectrum synthesizer 50 is designed to reconstruct replicas of the spectra of both voiced nasal and voiced nonnasal sounds, 75

while unvoiced spectrum synthesizer 51 is designed to reconstruct replicas of the spectra of unvoiced sounds.

Within synthesizer 50, there is provided a pulse generator 501 that delivers a succession of brief pulses of uniform amplitude whose repetition frequency is controlled by a pitch control signal of the type derived by detector 109 of FIG. 2C. If desired, a relaxation oscillator of conventional design may be used as a pulse generator 501. From pulse generator 501 the succession of pulses is passed to modulator 502, where the amplitudes of the incoming pulses are adjusted in response to a voiced amplitude control signal from detector 109 applied to the control terminal of modulator 502. The amplitude adjusted pulses appearing at the output terminal of modulator 502 are applied to a cascade of four electronically controllable resonance circuits 503 through 506 and one electronically controllable antiresonance circuit 507. Schematic diagrams of resonance and antiresonance circuits which may be utilized are shown in FIGS. 3A and 4A, respectively.

For the production of voiced nonnasal sounds, only the first three resonance circuits 503, 504, 505 are utilized, that is, voiced nonnasal control signals of the type developed by signal generator 105a are used to control resonance circuits 503, 504, 505 to reconstruct a replica of the spectrum of a voiced nonnasal sound. For the production of voiced nasal sounds, however, resonance circuit 506 and antiresonance circuit 507 are employed in addition to circuit 503 through 505, with a set of voiced nasal control signals of the type developed by signal generator 105b being furnished to control circuits 503 through 507. Thus synthesizer 50 reconstructs replicas of both voiced nasal and voiced nonnasal sounds, the particular spectrum depending upon the set of control signals which is applied to circuits 503 through 507.

The spectra of unvoiced sounds are reconstructed by synthesizer 51, which comprises a noise generator 511 for generating a noise voltage followed by a modulator 512. The bandwidth of the noise voltage depends upon the application of the unvoiced synthesizer, since as previously mentioned in connection with locator 11, it is necessary for unvoiced synthesizer in locator 11 to have a wide band noise voltage, whereas at the receiver station, the bandwidth of the noise voltage may be on the order of the bandwidth of the incoming speech wave from source 10. The amplitude of the noise voltage is adjusted in modulator 512 in response to an unvoiced amplitude control signal supplied, for example, by an unvoiced amplitude detector of the type illustrated by element 110 in FIG. 1. From modulator 512 the amplitude adjusted noise voltage is applied to a cascade of two electronically controllable resonant circuits 513 and 514 and one electronically controllable antiresonant circuit 515, examples of suitable resonant and antiresonant circuits being shown in FIGS. 3A and 4A, respectively. Reconstruction of a replica of an unvoiced sound spectrum is controlled by a set of unvoiced control signals of the type developed by signal generator 105c shown in FIG. 7, with the two resonance control signals being applied to the respective control terminals of circuits 513 and 514, and the one antiresonance control signal being applied to the control terminal of circuit 515.

Synthesizers 50 and 51 are connected to the input terminals of adder 52 so that the two synthesizers may operate simultaneously, if desired, for example, to reconstruct so-called voiced fricative sounds. The output terminal of adder 52 is connected to reproducer 53, which may be a conventional loudspeaker, in order to convert the reconstructed spectra into audible speech sounds.

Turning now to FIG. 3A, this is a schematic drawing of a resonant circuit comprising a resistive element R, an inductive element L, and a capacitive element C, all connected in series. The frequency of resonance of this circuit may be varied by changing the capacitance of capacitive element C, where if desired, C may be a well-

known Miller integrator circuit. The change in capacitance may be effected by a resonance control signal of the type previously described, and FIG. 3B illustrates an amplitude response curve for the circuit of FIG. 3A, in which the variable resonant frequency is denoted ω_0 .

The circuit of FIG. 3A permits the half-power bandwidth of the resonant frequency to be varied, if desired, by utilizing a suitable control signal to change the resistance of resistive element R. However, in the absence of a bandwidth control signal, the resistance of resistive element R may be set to a preassigned value corresponding to any desired resonance bandwidth appropriate to one or more classes of human speech sounds. For example, element R of resonance circuits 503, 504, 505, and 506 of synthesizer 50 in FIG. 5 may be adjusted to produce half-power bandwidths on the order of 50, 50, 75 and 200 cycles per second, respectively, while element R of resonant circuits 513 and 514 of synthesizer 51 may be adjusted to produce half-power bandwidths of about 500 cycles per second.

FIG. 4A is a schematic diagram of an antiresonant circuit suitable for use in voiced and unvoiced synthesizers 50 and 51 shown in FIG. 5. The voltage e_2 across the series connected inductive element 42, resistive element 43 and capacitive element 44 is given by

$$e_2 = i \left(j\omega L + R + \frac{1}{j\omega C} \right) \quad (1)$$

or

$$\frac{e_2}{i} = \left(j\omega L + R + \frac{1}{j\omega C} \right) \quad (2)$$

$$= \frac{L}{j\omega} \left(-\omega^2 + j \left(\frac{R}{L} \right) \omega + LC \right) \quad (3)$$

where $j = \sqrt{-1}$, ω denotes frequency, L is the inductance of element 42, R is the resistance of element 43, and C is the capacitance of element 44.

In order to produce an antiresonance at a desired frequency in the incoming signal to the circuit, denoted e_1 , it is necessary for the amplitude response curve for the circuit, as specified by the ratio e_2/e_1 , to have the shape illustrated in FIG. 4B, that is, at zero frequency, $\omega=0$, the gain must be unity, at the desired antiresonant frequency ω_0' , the gain must approach zero, and at frequencies higher than ω_0' the gain must increase with increasing frequency. These conditions are satisfied if the current, i , is made equal to

$$i = j\omega C e_1 \quad (4)$$

then Equation 3 becomes

$$\frac{e_2}{e_1} = LC \left(-\omega^2 + j \left(\frac{R}{L} \right) \omega + LC \right) \quad (5)$$

It is therefore evident from Equations 4 and 5 that the proper amplitude response is realized by passing the incoming signal e_1 , through a differentiator 40, which may be of any well-known construction, followed by an amplifier 41 having a high output impedance and delivering an output current that is proportional to its input voltage e_{in} according to the relationship

$$i = C e_{in} \quad (6)$$

where, by the action of differentiator 40,

$$e_{in} = j\omega e_1 \quad (7)$$

Amplifier 41 may be realized by a conventional remote cutoff pentode with a variable transconductance characteristic.

Variation of the antiresonant frequency ω_0' of the circuit shown in FIG. 4A may be achieved by changing the capacitance of element 44, as specified by Equation 5. A suitable embodiment for element 44 is a conventional Miller integrator circuit actuated by an antiresonance control signal of the type previously described. How-

ever, in order to preserve the amplitude response specified by Equation 5 for variations in the capacitance C of element 44, the factor C in Equation 6 requires that the output current delivered by amplifier 41 also be varied. This is accomplished by simultaneously applying the antiresonance control signal to amplifier 41 as well as to element 44.

The half-power bandwidth of the antiresonant frequency of the circuit of FIG. 4A may also be varied, if desired, by changing the resistance of element R in response to an appropriate control signal. However, in the absence of a bandwidth control signal, the resistance of element R may be adjusted to any desired value corresponding to the half-power bandwidth of one or more antiresonances of human speech sounds. For example, element R of antiresonant circuit 507 of synthesizer 50 in FIG. 5 may be adjusted to produce a half-power bandwidth of approximately 200 cycles per second, while element R of antiresonant circuit 515 in synthesizer 51 may be adjusted to produce a half-power bandwidth of approximately 1,000 cycles per second.

Although this invention has been described in terms of speech communications systems of the type shown in FIG. 1, it is to be understood that applications of the principles of this invention are not limited to such systems, but include the fields of automatic speech recognition, speech processing, and automatic message recording and reproduction. In addition, it is to be understood that the above-described embodiments are merely illustrative of the numerous arrangements which may be devised for the principles of this invention by those skilled in the art without departing from the spirit and scope of the invention.

What is claimed is:

1. A speech transmission system that comprises a transmitter station including
 - a source of an incoming speech wave characterized by a succession of different speech sounds including voiced nonnasal sounds, voiced nasal sounds, and unvoiced sounds,
 - means supplied with said speech wave for deriving from said succession of different sounds a corresponding succession of groups of reduced bandwidth control signals representative of selected resonances and antiresonances in the spectra of said different sounds of said speech wave, and
 - means supplied with said speech wave for obtaining a set of control signals representative of the frequencies and the amplitudes of the fundamental components of said voiced nonnasal and voiced nasal sounds of said speech wave and the energy of said unvoiced sounds of said speech wave,
 - means for transmitting said corresponding succession of groups of control signals and said set of control signals to a receiver station, and
 - at said receiver station,
 - means for synthesizing from said corresponding succession of groups of control signals and said set of control signals a succession of spectra which is a replica of the spectra of said succession of different sounds of said speech wave.
2. A speech transmission system that comprises a transmitter station including
 - a source of an incoming speech wave characterized by a succession of different speech sounds including voiced nonnasal sounds, voiced nasal sounds, and unvoiced sounds,
 - a resonance and antiresonance locator supplied with said speech wave for deriving from said succession of different sounds a corresponding succession of sets of control signals respectively representative of selected resonances and antiresonances in the spectrum of each of said different sounds of said speech wave, wherein a first one of said sets of control signals is representative of a plurality of selected resonances

of voiced nonnasal sound spectra, a second one of said sets of control signals is representative of a plurality of selected resonances and a single selected antiresonance of voiced nasal sound spectra, and a third one of said sets of control signals is representative of a plurality of selected resonances and a single selected antiresonance of unvoiced sound spectra,

a first detector for deriving from said speech wave a pitch control signal indicative of the frequency of the fundamental component of said voiced nonnasal sounds and said voiced nasal sounds of said speech wave and a voiced amplitude control signal indicative of the amplitude of the fundamental component of said voiced nonnasal sounds and said voiced nasal sounds of said speech wave,

a second detector for deriving from said speech wave an unvoiced amplitude control signal indicative of the energy of said unvoiced sounds of said speech wave,

means for transmitting to a receiver station said succession of sets of control signals, said pitch control signal, said voiced amplitude control signal, and said unvoiced amplitude control signal, and at said receiver station,

first synthesizing means supplied with said first and second sets of control signals, said pitch control signal, and said voiced amplitude control signal for reconstructing the spectra of said voiced nonnasal and said voiced nasal sounds of said speech wave,

second synthesizing means supplied with said third set of control signals and said unvoiced amplitude control signal for reconstructing the spectra of said unvoiced sounds of said speech wave, and

means for combining said reconstructed spectra to form a succession of spectra that is a replica of the spectra of said succession of different sounds of said incoming speech wave.

3. Apparatus for locating resonances and antiresonances in the spectra of a succession of different classes of speech sounds which comprises

a source of an incoming speech wave characterized by a succession of different speech sounds including voiced nonnasal sounds, voiced nasal sounds, and unvoiced sounds,

a source of a first control signal representative of the amplitudes of the fundamental frequency components of voiced nonnasal and voiced nasal sound portions of said speech wave,

a source of a second control signal representative of the frequency of the fundamental frequency components of voiced nonnasal and nasal sound portions of said speech wave,

a source of a third control signal representative of the energy of unvoiced sound portions of said speech wave,

a source of a repetitive cycle control signal having a predetermined repetition rate,

three signal generators in one-to-one correspondence with said voiced nonnasal, voiced nasal, and unvoiced speech sounds, respectively, wherein each of said signal generators, in response to said cycle control signal, produces a set of resonance and antiresonance signals whose values collectively vary to represent substantially all possible combinations of locations of selected resonances and antiresonances of the corresponding one of said classes of speech sounds during each cycle of said cycle control signal, said first signal generator producing a set of resonance signals which correspond to a plurality of selected resonances of voiced nonnasal sounds, said second signal generator producing a set of resonance and antiresonance signals which correspond to a plurality of selected resonances and a single selected antiresonance of voiced nasal sounds, and said third

signal generator producing a set of resonance and antiresonance signals which correspond to a plurality of selected resonances and a single selected antiresonance of unvoiced sounds,

a switching circuit provided with three input terminals and three output terminals in one-to-one correspondence with said plurality of signal generators for delivering said cycle control signal to a selected one of said signal generators according to the class of speech sound present in said speech wave,

means for applying said speech wave to a selected one of the input terminals of said switching circuit,

means for applying said first control signal to another of the input terminals of said switching circuit,

means for applying said cycle control signal to the third of the input terminals of said switching circuit,

means for connecting each output terminal of said switching circuit to one of said signal generators,

means supplied with said first, second, and third control signals for synthesizing from the selected set of signals produced by said selected one of said signal generators an artificial signal,

first analyzing means for deriving from said artificial signal a first group of signals representative of the amplitudes of the individual frequency components of the spectrum of said artificial signal,

second analyzing means for deriving from said speech wave a second group of signals representative of the amplitudes of the individual frequency components of the spectrum of said speech wave,

comparing means in circuit relation with said first and second analyzing means for obtaining from said first and second groups of amplitude signals an error signal indicative of a preassigned measure of the difference between said artificial spectrum and said speech spectrum, and

means responsive to said error signal for sampling the values of said selected set of signals which correspond to the minimum magnitude of said error signal during each cycle of said cycle control signal.

4. Apparatus as defined in claim 3 wherein said switching circuit comprises

a first relay means provided with an input terminal, a control terminal, and first and second output terminals,

a second relay means provided with an input terminal, a control terminal, and first and second output terminals,

means for applying said first control signal to the input terminal of said first relay means,

means for connecting the first output terminal of said first relay means to the input terminal of said third signal generator,

means for connecting the second output terminal of said first relay means to the input terminal of said second relay means,

means for connecting the first output terminal of said second relay means to the input terminal of said second signal generator,

means for connecting the second output terminal of said second relay means to the input terminal of said first signal generator,

means for distinguishing between voiced nasal sounds and voiced nonnasal sounds which includes

a first signal path provided with an input terminal, an output terminal, and comprising a first bandpass filter with a pass band extending from about zero to 1,500 cycles per second,

a first rectifier, and a first low pass filter connected in series,

a second signal path provided with an input terminal, an output terminal, and comprising a second bandpass filter with a pass band extending from about 1,500 to 3,000 cycles per second, a second rectifier, and a second low pass filter connected in series,

15

means for simultaneously applying said speech wave to the input terminals of said first and second signal paths,
 subtracting means provided with a minuend terminal, a subtrahend terminal, and an output terminal, 5
 means for connecting the output terminal of said first signal path to the minuend terminal of said subtracting means,
 means for connecting the output terminal of said second signal path to the subtrahend terminal of said subtracting means, 10
 a threshold detector provided with an input terminal, an output terminal, and a threshold corresponding to the smallest expected difference in energy between the group of frequency components lying in the frequency range under about 1,500 cycles per second, and the group of frequency components lying within the frequency range that extends from about 1,500 to 3,000 cycles per second, and 15
 means for connecting the output terminal of said subtracting means to the input terminal of said threshold detector, and 20
 means for connecting the output terminal of said threshold detector to the control terminal of said second relay means. 25

5. Apparatus for deriving from a speech wave control signals indicative of the frequency and the amplitude of the fundamental speech frequency component which comprises

a source of an incoming speech wave, 30
 filter means supplied with said speech wave for passing the fundamental component of said speech wave, wherein said filter means is provided with a

$$\frac{1}{f^2}$$

characteristic in the interval between about 70 and 250 cycles per second, where f denotes frequency, 40
 a first subpath comprising a rectifier and a low pass filter connected in series for rectifying and averaging the fundamental component passed by said filter means,
 a second subpath connected in parallel with said first subpath comprising frequency measuring means for obtaining a first control signal having a magnitude representative of the frequency of the fundamental component passed by said filter means, 45
 means supplied with said first control signal for squaring the magnitude of said first control signal,
 multiplier means in circuit relationship with said first subpath and said squaring means for multiplying said rectified and averaged fundamental component by the squared magnitude of said first control signal to produce a second control signal indicative of the magnitude of the fundamental component passed by said filter means. 50

6. A speech synthesizer for reconstructing a replica of an incoming speech wave characterized by a succession of different sounds including voiced nonnasal sounds, voiced nasal sounds, and unvoiced sounds comprising 60

a source of a pitch control signal representative of the frequency of the fundamental component of each of said voiced nonnasal and voiced nasal speech sounds,
 a source of a voiced amplitude control signal representative of the amplitude of the fundamental component of said voiced nonnasal sounds and said voiced nasal speech sounds, 65
 a source of a first sequence of sets of control signals representative of selected resonance and antiresonance locations in the amplitude spectra of said corresponding succession of voiced nonnasal and voiced nasal speech sounds, 70
 first synthesizer means responsive to said pitch con- 75

16

trol signal, said voiced amplitude control signal, and said first sequence of sets of control signals for reconstructing the amplitude spectra of said succession of voiced nonnasal and voiced nasal speech sounds, 5
 a source of an unvoiced amplitude control signal representative of the energy of unvoiced speech sounds, a source of a second set of control signals representative of selected resonance and antiresonance locations in the amplitude spectra of each of said unvoiced sounds in said succession of different sounds in said speech wave, 10

second synthesizer means responsive to said unvoiced amplitude control signal and said second set of resonance and antiresonance control signals for reconstructing the amplitude spectra of unvoiced speech sounds in said succession of different sounds in said speech wave, 15

adding means for combining the spectra reconstructed by said first and second synthesizer means to form a succession of spectra that is a replica of the spectra of said succession of voiced nonnasal sounds, voiced nasal sounds, and unvoiced sounds in said speech wave, and 20

reproducing means supplied with the succession of spectra from said adding means for converting said succession of spectra into a corresponding succession of audible speech sounds. 25

7. Apparatus as defined in claim 6 wherein said first synthesizer means for reconstructing the amplitude spectra of said succession of voiced nonnasal and voiced nasal speech sounds comprises 30

means responsive to said pitch control signal for generating a train of uniform amplitude pulses with a repetition rate equal to the frequency of said fundamental component, 35

means responsive to said voiced amplitude control signal and supplied with said train of pulses for adjusting the amplitudes of said pulses to correspond to the amplitude of said fundamental component, and 40

a cascade of four resonant circuits and one antiresonant circuit supplied with said train of amplitude adjusted pulses and responsive to said first sequence of sets of control signals for reconstructing a succession of voiced nonnasal amplitude spectra and voiced nasal amplitude spectra corresponding to said succession of voiced nonnasal sounds and voiced nasal sounds in said incoming speech wave. 45

8. Apparatus as defined in claim 6 wherein said second synthesizer means for reconstructing the amplitude spectra of said succession of unvoiced speech sounds comprises 50

means for generating a noise signal, means responsive to said unvoiced amplitude control signal and supplied with said noise signal for adjusting the amplitude of said noise signal to correspond to the energy of said unvoiced speech sounds, and 55

a cascade of two resonant circuits and one antiresonant circuit supplied with said amplitude adjusted noise signal and responsive to said second set of control signals for reconstructing a succession of unvoiced amplitude spectra corresponding to said succession of unvoiced speech sounds in said incoming speech wave. 60

9. Apparatus for constructing an amplitude spectrum with an antiresonance that can be varied in location which comprises 65

a source of an incoming signal, a source of an antiresonance control signal having a magnitude which is representative of a relatively wide range of antiresonance locations, 70

a differentiator provided with an input terminal and an output terminal for developing from said incoming signal an output signal representative of the derivative of said incoming signal, 75

17

means for applying said incoming signal to the input
terminal of said differentiator,
a variable transconductance amplifier provided with
an input terminal, an output terminal and a control
terminal, 5
an inductive element, a resistive element, and a variable
capacitive element connected in series,
means for connecting the output terminal of said amplifier
to said inductive element, and
means for applying said antiresonance control signal 10
to the control terminal of said amplifier and to vary
the capacitance of said variable capacitive element,

18

whereby the output signal developed across said series
connected inductive, resistive, and variable capacitive
elements contains an antiresonance whose location
is determined by the magnitude of said antiresonance
control signal.

References Cited by the Examiner

UNITED STATES PATENTS

2,817,707 12/57 Weibel ----- 179—1

DAVID G. REDINBAUGH, *Primary Examiner.*