

(19) World Intellectual Property Organization
International Bureau



(43) International Publication Date
14 February 2008 (14.02.2008)

PCT

(10) International Publication Number
WO 2008/018653 A1

(51) International Patent Classification:
G10L 13/02 (2006.01)

(74) Agents: KWON, Oh-Sig et al.; 4F, Joeeun Leaderstel, 921
Dunsan-dong Seo-gu, Daejeon 302-120 (KR).

(21) International Application Number:
PCT/KR2006/004478

(22) International Filing Date: 31 October 2006 (31.10.2006)

(25) Filing Language: Korean

(26) Publication Language: English

(30) Priority Data:
10-2006-0075140 9 August 2006 (09.08.2006) KR

(71) Applicant (for all designated States except US): KOREA
ADVANCED INSTITUTE OF SCIENCE AND TECH-
NOLOGY [KR/KR]; 373-1 Guseong-dong, Yuseong-gu,
Daejeon 305-701 (KR).

(81) Designated States (unless otherwise indicated, for every
kind of national protection available): AE, AG, AL, AM,
AT, AU, AZ, BA, BB, BG, BR, BW, BY, BZ, CA, CH, CN,
CO, CR, CU, CZ, DE, DK, DM, DZ, EC, EE, EG, ES, FI,
GB, GD, GE, GH, GM, GT, HN, HR, HU, ID, IL, IN, IS,
JP, KE, KG, KM, KN, KP, KZ, LA, LC, LK, LR, LS, LT,
LU, LV, LY, MA, MD, MG, MK, MN, MW, MX, MY, MZ,
NA, NG, NI, NO, NZ, OM, PG, PH, PL, PT, RO, RS, RU,
SC, SD, SE, SG, SK, SL, SM, SV, SY, TJ, TM, TN, TR,
TT, TZ, UA, UG, US, UZ, VC, VN, ZA, ZM, ZW.

(84) Designated States (unless otherwise indicated, for every
kind of regional protection available): ARIPO (BW, GH,
GM, KE, LS, MW, MZ, NA, SD, SL, SZ, TZ, UG, ZM,
ZW), Eurasian (AM, AZ, BY, KG, KZ, MD, RU, TJ, TM),
European (AT, BE, BG, CH, CY, CZ, DE, DK, EE, ES, FI,
FR, GB, GR, HU, IE, IS, IT, LT, LU, LV, MC, NL, PL, PT,
RO, SE, SI, SK, TR), OAPI (BF, BJ, CF, CG, CI, CM, GA,
GN, GQ, GW, ML, MR, NE, SN, TD, TG).

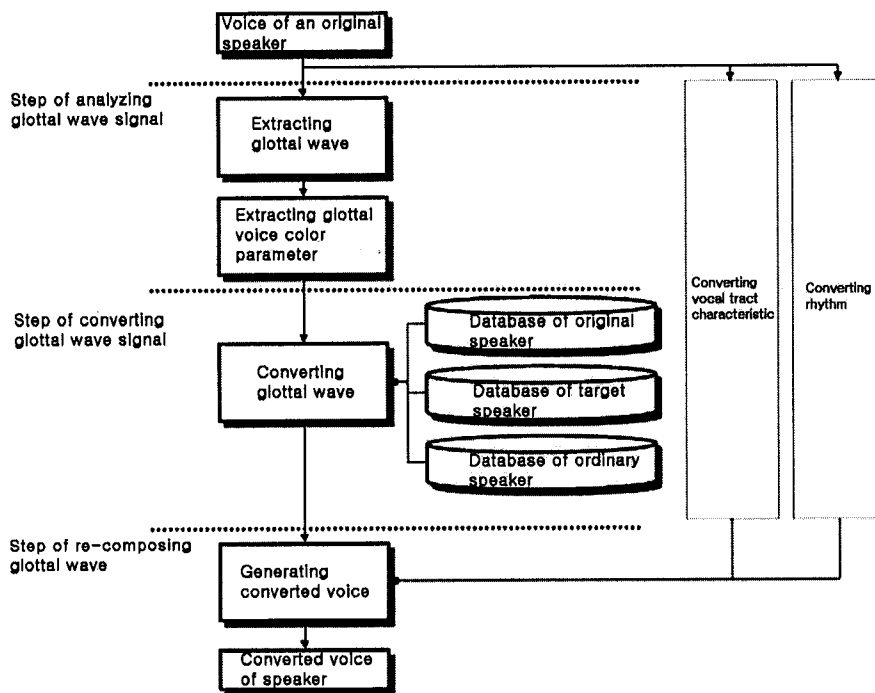
(72) Inventors; and

(75) Inventors/Applicants (for US only): OH, Yung-Hwan
[KR/KR]; 110-1206 Hanwool APT., Sinseong-dong,
Yuseong-gu, Daejeon 305-707 (KR). BAE, Jae-Hyun
[KR/KR]; 205 Ho, 111-14, Eoeun-dong, Yuseong-gu,
Daejeon 305-807 (KR).

Published:

— with international search report

(54) Title: VOICE COLOR CONVERSION SYSTEM USING GLOTTAL WAVEFORM



(57) Abstract: Disclosed is a speaker conversion system which can express various voice colors, more particularly, to voice color conversion method and system which can generate converted voice having various voice colors according to an utterance situation by converting the glottal wave.

WO 2008/018653 A1

VOICE COLOR CONVERSION SYSTEM USING GLOTTAL WAVEFORM

【Technical Field】

5 The present invention relates to a speaker conversion system which can express various voice colors, more particularly, to voice color conversion method and system which can generate converted voice having various voice colors according to an utterance situation by converting the glottal
10 wave.

【Background Art】

FIG. 1 shows a process of generating a voice upon speaking. When a man speaks, air from a bronchial tube passes
15 through the vocal cords and thus glottal wave is formed. At this time, since the man pronounces while breathing out his/her breath, noise (aspirated sound) due to out-breathing is contained in the glottal wave. While the glottal wave passes through a vocal tract, an articulating phenomenon is
20 occurred. Then, the glottal wave is radiated into the air and thus generates a voice. The present invention relates to the conversion of glottal wave generated when a man speaks and grows out of a fact that a shape of glottal wave is changed according to environmental status, feelings and the like when
25 a man speaks and thus the voice having various voice colors

can be generated.

As a conventional method of mimicking a voice of a specific person (hereinafter, called "target speaker"), there are a mimicking method employing an expert dubbing artist and
5 a mimicking method using a computer.

In case of the mimicking method employing an expert dubbing artist, it is possible to mimic a prosodic characteristic with respect to a certain part of the voice, however, it is difficult to express various and natural voice
10 colors.

And in case of the method using a computer, which can convert a voice of a certain speaker (hereinafter, called "original speaker") into the target speaker by using the computer, there are mimicking methods using Hidden Markov
15 Model (HMM), Gaussian Mixture Model (GMM) and Neural Network.

In the conventional methods using the HMM, the GMM and the Neural Network, firstly, vocal tract characteristic parameters of voice such as LPC (Linear Prediction Coefficient), LSP (Line Spectral Pair), MFCC (Mel-Frequency
20 Cepstral Coefficient) and HNM (Harmonic and Noise Model) characteristics are extracted from the voices of the original speaker and the target speaker, and the HMM or GMM is trained by using the characteristic parameters of each speaker, and then a conversion function between the trained models is
25 calculated so that the vocal tract characteristic of the

original speaker is converted into that of the target speaker. In case of the prosody, the prosody of the target speaker is modeled and then applied to the converted voice. In a method of mimicking the prosody of the target speaker, after forming
5 a pitch histogram of the original speaker and the target speaker, an excitation signal matched with the histogram is used. In case of the glottal wave, the excitation signal in which basic frequency information of the target speaker is excluded is applied to the converted voice.

10 In the conventional speaker conversion method, it is possible to convert the prosody and the vocal tract characteristic. However, since the excitation signal of the target speaker is used in the vocal cord characteristic, there is a problem that the voice color contained in the given
15 excitation signal of the target speaker has to be used as it is.

Therefore, in the conventional method it is difficult to exactly express various voice colors according to situations, it is impossible to reflect the change of voice color
20 according to feelings, environments, and the context when the original speaker speaks, and it is possible only to convert the voice of the original speaker into a predetermined voice color.

25 **【Disclosure】**

【Technical Problem】

An object of the present invention is to provide a method of converting voice color, which can generate a voice of a certain speaker without employment of an expert dubbing artist who can mimic the voice and also which can generate voices having various voice color at anywhere and any time according to feelings of the certain person.

Another object of the present invention is to provide a method of converting voice color, which can generate voices having various voice colors according to situations and conditions

【Technical Solution】

According to the present invention, there is provided a voice color conversion method for converting a speaker, comprising the steps of analyzing a glottal wave signal of a voice of an original speaker; converting the analyzed glottal wave signal; and re-composing the converted glottal wave signal.

And the step of analyzing the glottal wave signal comprises the steps of extracting a glottal wave from the voice of the original speaker; and extracting voice color characteristic parameters of the extracted glottal wave.

Further, the step of converting the glottal wave signal comprises the steps of collecting glottal wave characteristic

parameters for various voice colors from the voice of the original speaker and creating a voice color database of the original speaker; and collecting glottal wave characteristic parameters for various voice colors from a voice of a target speaker and creating a voice color database of the target speaker. The step of converting the glottal wave signal further comprises the steps of collecting glottal wave characteristic parameters for various voice colors from voices of various general speakers and creating a voice color database of the general speakers. The step of converting the glottal wave signal further comprises the steps of creating a corresponding relationship among the voice color databases of the original speaker, the target speaker and the general speakers on the basis of a quantity of change of the voice color characteristic parameters stored in each voice color database.

Further, the characteristic parameters comprises an OQ section in which vocal cords are opened, a CQ section in which the vocal cords are closed, EE (Effective Excitation) which is an effective excitation value and T_0 which is a period of signal. The quantity of change of the characteristic parameters is an NAQ parameter which is defined as follows:

$$NAQ = \frac{E_0}{T_0 EE}$$

In addition, a voice color conversion system for converting a speaker, comprising glottal wave extracting means for extracting a glottal wave from a voice of an original speaker; voice color parameter extracting means for extracting
5 voice color parameters from the extracted glottal wave; glottal wave converting means for converting the glottal wave using the voice color parameters; and converted voice generating means for generating a converted voice using the converted glottal wave.

10 Preferably, the glottal wave extracting means comprises an A/D converter and an input buffer, and the voice color parameter extracting means and the glottal wave converting means comprises a command storing unit and a main controller, and the converted voice generating means comprises a D/A
15 converter and an output buffer. Further, The system further comprises a storing unit in which voice color databases of an original speaker, a target speaker and general speakers are stored.

20 **【Advantageous Effects】**

According to the voice color conversion method of the present invention, it is possible to convert a voice of an original speaker in to a voice having various voice color of a target speaker according to situations, feelings and context.

25 Furthermore, it also is possible that the speaker

conversion system of the present invention can facilely substitute for an expert dubbing artist who can mimic a voice of a target speaker. Further, since the speaker conversion system can optionally generate a virtual voice by combination
5 of the voice color database of the general speaker, it can be used to make voices of various virtual characters.

【Description of Drawings】

The above and other objects, features and advantages of
10 the present invention will become apparent from the following description of preferred embodiments given in conjunction with the accompanying drawings, in which:

Fig. 1 is a view showing a mechanism for generating a voice when a mean speaks;

15 Fig. 2 is a view showing a process of converting glottal wave to express various voice colors according to the present invention;

Fig. 3 is a view showing a glottal wave obtained by using an LF (Liljencrants-Fant) model and a derivative glottal wave
20 thereof according to the present invention;

Fig. 4 is a view showing a derivative glottal wave of a KLGLOT88 model according to the present invention;

Fig. 5 is a view showing characteristic parameters of glottal wave for converting the voice color according to the
25 present invention;

Fig. 6 is a view showing a manner of storing voice color data in voice color databases of an original speaker, a target speaker and a general speaker according to the present invention;

5 Fig. 7 is a graph showing distribution of NAQ parameters for various voice colors according to the present invention;

Fig. 8 is a view showing a change of vocal glottal waves having various voice colors according to the present invention;

10 Fig. 9 is a view showing a process in a step of converting the glottal wave according to the present invention; and

Fig. 10 is a block diagram explaining a system for converting glottal wave signal according to the present
15 invention.

【Best Mode】

Hereinafter, the embodiments of the present invention will be described in detail with reference to accompanying
20 drawings.

Fig. 2 is a view showing a process of converting glottal wave to express various voice colors according to the present invention. As shown in Fig. 2, a method of converting a voice according to the present invention includes a step of
25 analyzing glottal wave signal, a step of converting the

glottal wave signal and a step of re-composing the glottal wave signal.

The step of analyzing glottal wave signal includes a step of extracting glottal wave from voice signal of an input original speaker and a step of extracting parameters of various voice colors from the extracted glottal wave. In the step of converting the glottal wave signal, the parameters of voice colors extracted from the glottal wave are converted into glottal wave signals using data of voice color database. At this time, the voice color database includes voice color database of an original speaker, voice color database of a target speaker and voice color database of a general speaker.

The voice color database of the target speaker as a model representing native voice color of only the target speaker is a database in which characteristic parameters of the target speaker are collected. The voice color database of the general speaker is a database in which characteristic parameters of various voice colors extracted from general speakers are collected.

In the step of converting the glottal wave signal, the characteristic parameters of voice colors of the original speaker, which are extracted in the step of analyzing the glottal wave signal, are converted into the glottal wave having the voice color of the target speaker using the voice color characteristic parameters in which voice color models of

the target speaker and the general speaker weighted and combined with each other.

In the step of re-composing the glottal wave, the glottal wave having the voice color of the target speaker is reconstructed using vocal tract characteristic parameters by
5 using the HMM and the like and the glottal wave generated in the step of converting the glottal wave signal so as to generate a finally converted voice.

Meanwhile, in the step of analyzing the glottal wave, the
10 glottal wave is extracted from input signal which is converted into digital signal through a sampling process, and the characteristic parameters of voice colors are extracted from the glottal wave. When extracting the glottal wave in the step, the vocal tract characteristic is removed as much as possible
15 so as to remain only the glottal wave. The glottal wave may be extracted by inverse-filtering excitation signal using linear prediction algorithm of voice. To analyze the extracted the glottal wave, the glottal wave is differentiated along a time axis to obtain glottal derivative signal.

20 Fig. 3 is a view showing a glottal wave obtained by using an LF (Liljencrants-Fant) model and a derivative glottal wave thereof according to the present invention. The derivative glottal wave may be represented by using the LF model as follows:

$$\begin{aligned}
g(t) &= E_0 e^{\alpha t} \sin(\omega_{gt}) & , \quad 0 \leq t \leq T_e \\
&= -\frac{EE}{\epsilon T_a} \left[e^{-\epsilon(t-T_e)} - e^{-\epsilon(T_c-T_e)} \right], & T_e \leq t \leq T_c \leq T_0
\end{aligned} \tag{1}$$

wherein E_0 is a maximum amplitude of the glottal wave signal, EE is a maximum closing speed of the vocal cords as an effective Excitation value, T_e is a time period while the vocal cords are opened, T_c is a time period while the vocal cords are closed and T_0 is a period while the vocal cords are opened and closed.

To represent the glottal wave, a KLGLOT88 model may be also used. The equation is as follows:

$$\begin{aligned}
g(t) &= at^2 - bt^3 & , \text{ for } 0 < t < O_q T_0 \\
&= 0 & , \text{ for } O_q T_0 < t < T_0
\end{aligned}$$

(2)

wherein a and b are defined as a voicing amplitude (AV) and an equation of O_q , and the a and b may be represented as follows:

$$a = \frac{27AV}{4T_0 O_q^2}, \quad b = \frac{27AV}{4T_0 O_q^3}$$

(3)

wherein The AV is the same parameter as the E_0 of the LF model, O_q is also the same parameter as the OQ of the LF model

and the T_0 is a basic period of the vocal cords.

Fig. 4 is a view showing a derivative glottal wave of the KLGLOT88 model. As shown in Figs. 3 and 4, although the used models are respectively different, shapes of the glottal waves
5 are equal or similar to each other, because the same characteristic parameters are used.

Fig. 5 is a view showing characteristic parameters of glottal wave for converting the voice color according to the present invention. The characteristic parameters which is
10 necessary to express various voice colors includes OQ which indicates an opening section of the vocal cords, CQ which indicates a closing section of the vocal cords, EE (Effective Excitation) which is an effective excitation value and T_0 which is a period of signal. The characteristic parameters can be
15 obtained during a process of matching the derivative glottal wave with the LF model using the Figs. 3 and 4. And after modeling the LF derivative glottal wave, an error signal which indicates a difference between the characteristic parameters of the derivative glottal wave and the actual derivative
20 glottal wave can be obtained.

Fig. 6 is a view showing a manner of storing voice color data in the voice color databases of the original speaker, the target speaker and the general speaker according

the present invention. The voice color database of the
25 target speaker stores the characteristic parameters of glottal

wave which is analyzed in the signal analyzing step for the actual voice of the target speaker and the error signal for each voice color. The voice color database of the general speaker stores the characteristic parameters of glottal wave which is analyzed in the signal analyzing step for the actual voices of the ordinary people and the error signal for each voice color. A breathy voice, a pressed voice and a normal voice are representative of the voice colors which are expressed by the glottal wave. On the basis of the fact, the voice color databases of the target speaker and the general speaker stores the characteristic parameters of various voice colors in a storing manner shown in Fig. 6. In case of the breathy voice, the OQ and CQ sections are not divided clearly and the EE section has a flat shape. In case of the pressed voice, the EE section has a peak shape, and the CQ section is much longer than the OQ section. The normal voice has average OQ, CQ and EE sections. When generating the converted glottal wave, the glottal wave can have various voice colors by adjusting the characteristic parameters such as the OQ, the CQ, the EE, the E_0 and the like. The voice color database of the original speaker stores various voice colors of the original speaker, which are divided according to situations, feelings and the like. In the above three voice color databases, the various voice colors are correspondent to each other according to the divided conditions. The corresponding relationship is

set by the histogram of changes in the parameters. Therefore, the various voice colors of the original speaker can be converted into the voice colors of the target speaker by the voice color databases of the target speaker and the general speaker. The three voice color databases are constructed to have voice colors correspondent to each other according to voice color numbers.

Herein, to explain the corresponding relationship among the original speaker, the target speaker and the general speaker, it will be described that the speakers have common characteristics for various voice colors. To this end, the present invention introduces an NAQ parameter. The NAQ parameter is a combination of the above characteristic parameters of glottal wave and represented as follows:

$$NAQ = \frac{E_0}{T_0 EE}$$

(4)

That is, the characteristic parameters of glottal wave are extracted from the glottal wave shown in Fig. 5, and the NAQ parameter is obtained by using the Equation 4 to show the distribution thereof, whereby it is possible to grasp the corresponding relationship among the above characteristic parameters

Fig. 7 is a graph showing distribution of NAQ parameters

obtained by combining the characteristic parameters of glottal wave, wherein the glottal waves of breathy (Bre), normal (Nor) and pressed (Pre) voices pronounced by thirteen people are analyzed and then the distribution of NAQ parameters is
5 obtained. Referring to Fig. 7, the NAQ parameters are constantly distributed for a common voice color and thus it is the case that a quantity of change of the parameter values for each voice color is constant in each voice color. Therefore, it is possible to set the corresponding relationship for each
10 voice color.

Fig. 8 is a view showing a shape of derivative glottal waves for the three voice colors, in which the change of the characteristic parameters OQ, CQ and EE for each voice color is observed. Referring to Fig. 8, in case of the breathy voice,
15 the QC section and the CQ section are not clearly divided and the EE section has a gentle shape. However, in case of the pressed voice, the CQ section is relatively longer than the OQ section and the EE section has a peak shape.

In the step of converting the glottal wave, the
20 characteristic parameters of glottal wave and the error signal which have a desired voice color of the target speaker are extracted and then the parameter values are converted so as to generate the a glottal wave model having the desired voice color of the target speaker. If there is not the desired voice
25 color in the voice color database of the target speaker, the

characteristic parameters of the desired voice color and the characteristic parameters of the normal voice color are extracted from the voice color database of the general speaker, and the quantity of change between the two characteristic parameters are calculated. Then, the calculated quantity of change is applied to the glottal wave having the normal voice color of the target speaker, thereby generating a glottal wave model having the desired voice of the target speaker. The generated glottal wave model is combined with a basic frequency, a duration time and an energy parameter calculated in an existing prosody converting step to generate an actual glottal wave $g(t)$, and the generated glottal wave is passed through a filter $r(t)$ which has a radiation effect from the lips when a man speaks and then passed through a linear prediction filter $v(t)$ comprised of the vocal tract characteristic parameters calculated in an existing vocal tract characteristic converting step, thereby finally generating a desired voice of the target speaker.

Fig. 9 is a view showing a process in a step of converting the glottal wave according to the present invention, which is represented as follows:

$$s(t) = g(t) * r(t) * v(t)$$

(5)

wherein * indicates convolution

The filter $r(t)$ may be represented as follows:

$$r(t) = \frac{1}{1 - z^{-1}}$$

(6)

5 As described above, since the final voice which is generated through the signal conversion step can reproduce the voice of target speaker as it is, the present invention can provide various converted voices compared with the conventional speaker conversion system which can express only
10 one voice color.

Fig. 10 is a block diagram explaining a system for converting glottal wave signal according to the present invention. As shown in Fig. 10, the voice of speaker is input to an A/D converter and an input buffer to extract a glottal
15 wave, and a voice color parameter of the extracted glottal wave are extracted by a command storing unit and a main controller and converted into a glottal wave, and then the final voice is output through an D/A converter and an output buffer.

20

【Industrial Applicability】

According to the voice color conversion method of the present invention, it is possible to convert a voice of an original speaker in to a voice having various voice color of a

target speaker according to situations, feelings and context.

Further, the speaker conversion system of the present invention can be widely applied to animations, movies, plays, commercial films and the like where voices having various and
5 plentiful voice colors are needed.

Furthermore, it also is possible that the speaker conversion system of the present invention can facilely substitute for an expert dubbing artist who can mimic a voice of a target speaker. Further, since the speaker conversion
10 system can optionally generate a virtual voice by combination of the voice color database of the general speaker, it can be used to make voices of various virtual characters.

【CLAIMS】**【Claim 1】**

A voice color conversion method for converting a speaker,
comprising the steps of:

5 analyzing a glottal wave signal of a voice of an original
speaker;

converting the analyzed glottal wave signal; and

re-composing the converted glottal wave signal.

10 **【Claim 2】**

The method according to claim 1, wherein the step of
analyzing the glottal wave signal comprises the steps of:

extracting a glottal wave from the voice of the original
speaker; and

15 extracting voice color characteristic parameters of the
extracted glottal wave.

【Claim 3】

The method according to claim 1, wherein the step of
20 converting the glottal wave signal comprises the steps of:

collecting glottal wave characteristic parameters for
various voice colors from the voice of the original speaker
and creating a voice color database of the original speaker;
and

25 collecting glottal wave characteristic parameters for

various voice colors from a voice of a target speaker and creating a voice color database of the target speaker.

【Claim 4】

5 The method according to claim 3, wherein the step of converting the glottal wave signal further comprises the steps of collecting glottal wave characteristic parameters for various voice colors from voices of various general speakers and creating a voice color database of the general speakers.

【Claim 5】

10 The method according to claim 4, wherein the step of converting the glottal wave signal further comprises the steps of creating a corresponding relationship among the voice color
15 databases of the original speaker, the target speaker and the general speakers on the basis of a quantity of change of the voice color characteristic parameters stored in each voice color database.

【Claim 6】

20 The method according to claim 5, wherein the characteristic parameters comprises an OQ section in which vocal cords are opened, a CQ section in which the vocal cords are closed, EE (Effective Excitation) which is an effective
25 excitation value and T_0 which is a period of signal.

【Claim 7】

The method according to claim 5 or 6, wherein the quantity of change of the characteristic parameters is an NAQ
5 parameter which is defined as follows:

$$NAQ = \frac{E_0}{T_0 EE}$$

【Claim 8】

10 A voice color conversion system for converting a speaker, comprising:

glottal wave extracting means for extracting a glottal wave from a voice of an original speaker;

voice color parameter extracting means for extracting
15 voice color parameters from the extracted glottal wave;

glottal wave converting means for converting the glottal wave using the voice color parameters; and

converted voice generating means for generating a converted voice using the converted glottal wave.

20

【Claim 9】

The system according to claim 8, wherein the glottal wave extracting means comprises an A/D converter and an input buffer.

【Claim 10】

The system according to claim 8, wherein the voice color
parameter extracting means and the glottal wave converting
5 means comprises a command storing unit and a main controller.

【Claim 11】

The system according to claim 8, wherein the converted
voice generating means comprises a D/A converter and an output
10 buffer.

【Claim 12】

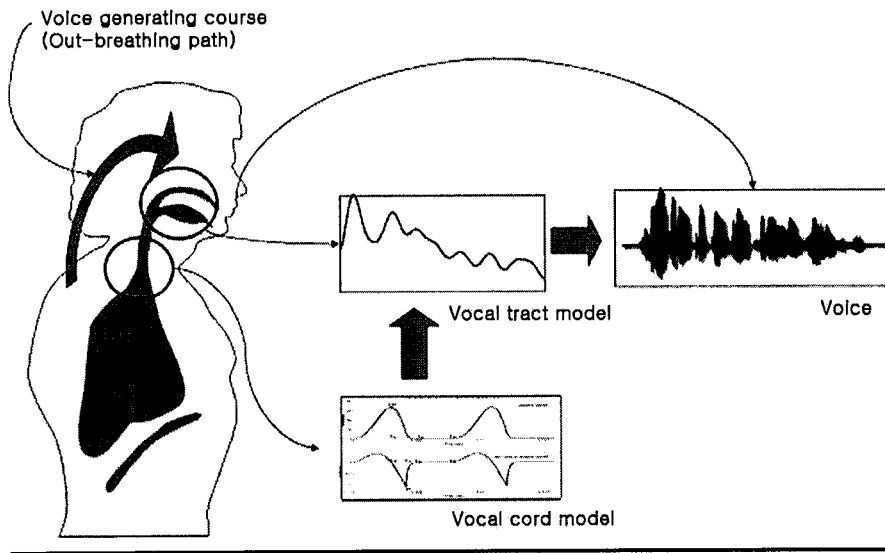
The system according to claim 8, further comprising a
storing unit in which voice color databases of an original
15 speaker, a target speaker and general speakers are stored.

20

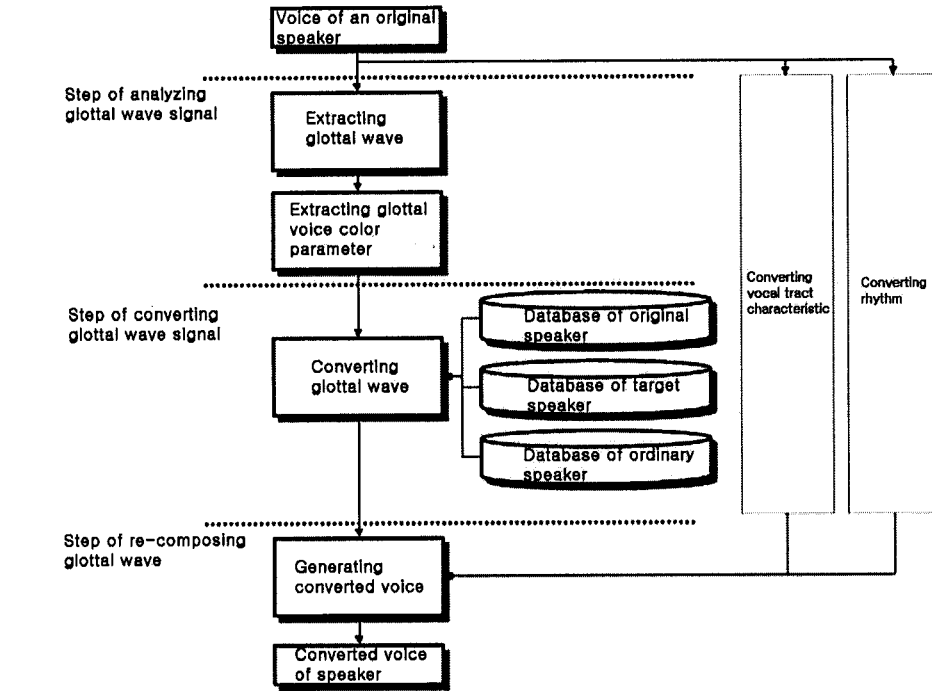
25

【DRAWINGS】

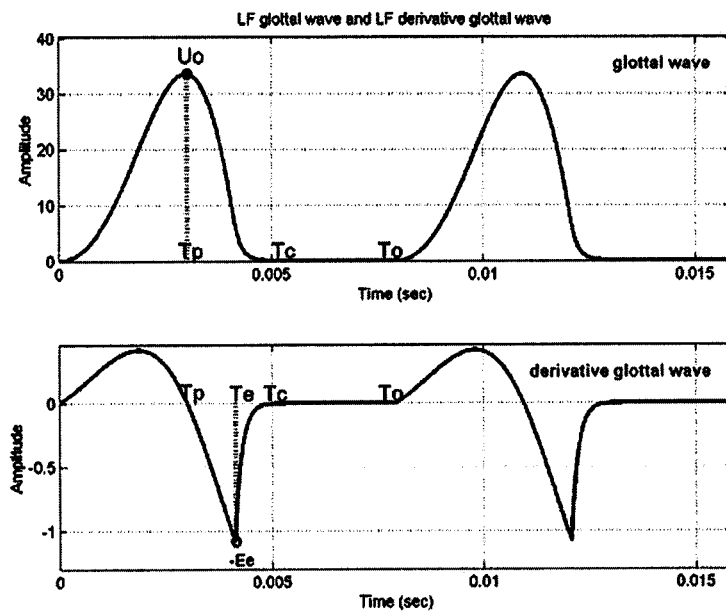
【Figure 1】



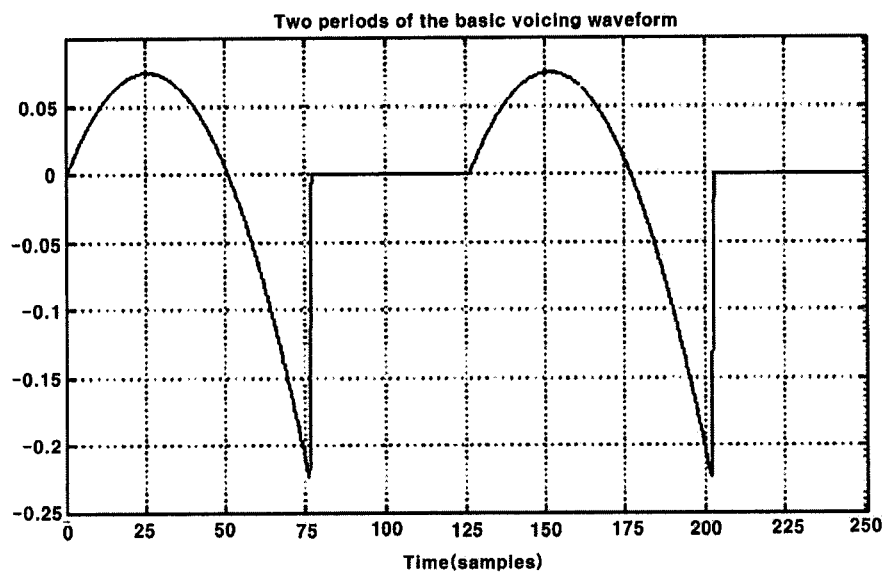
【Figure 2】



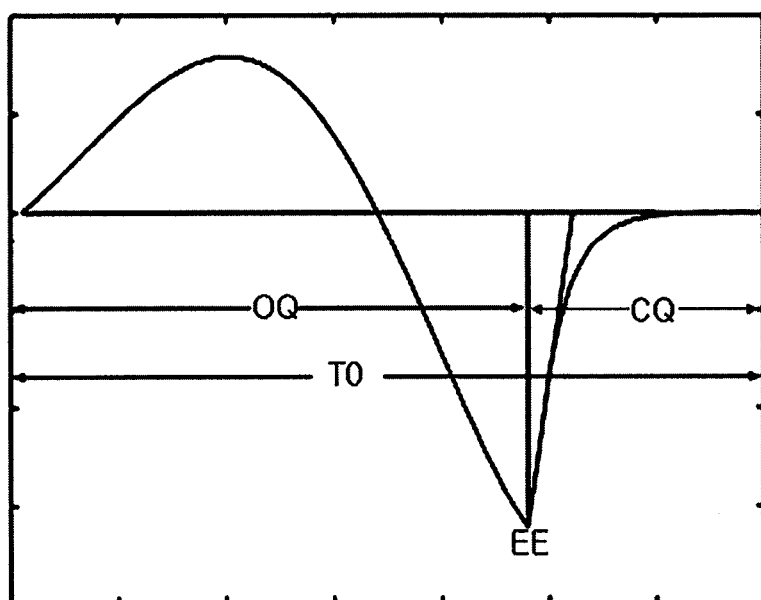
【Figure 3】



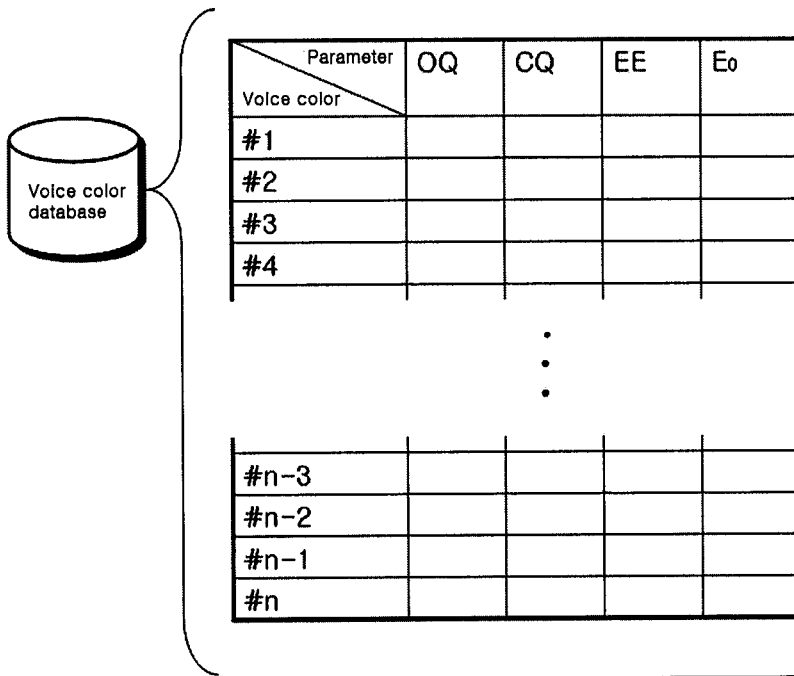
【Figure 4】



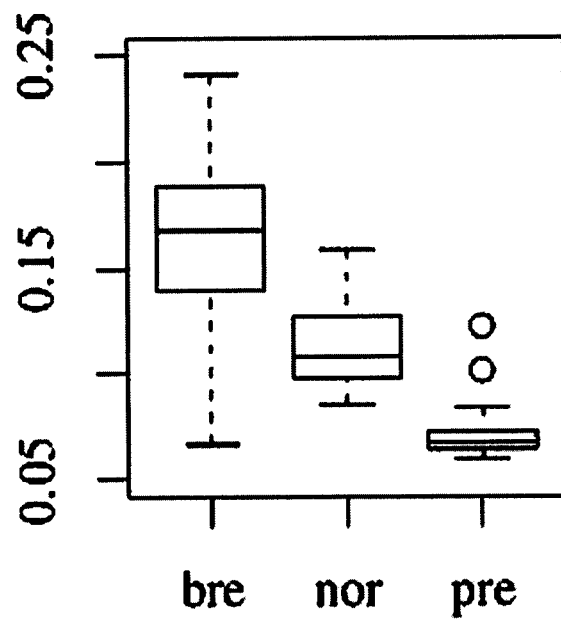
【Figure 5】



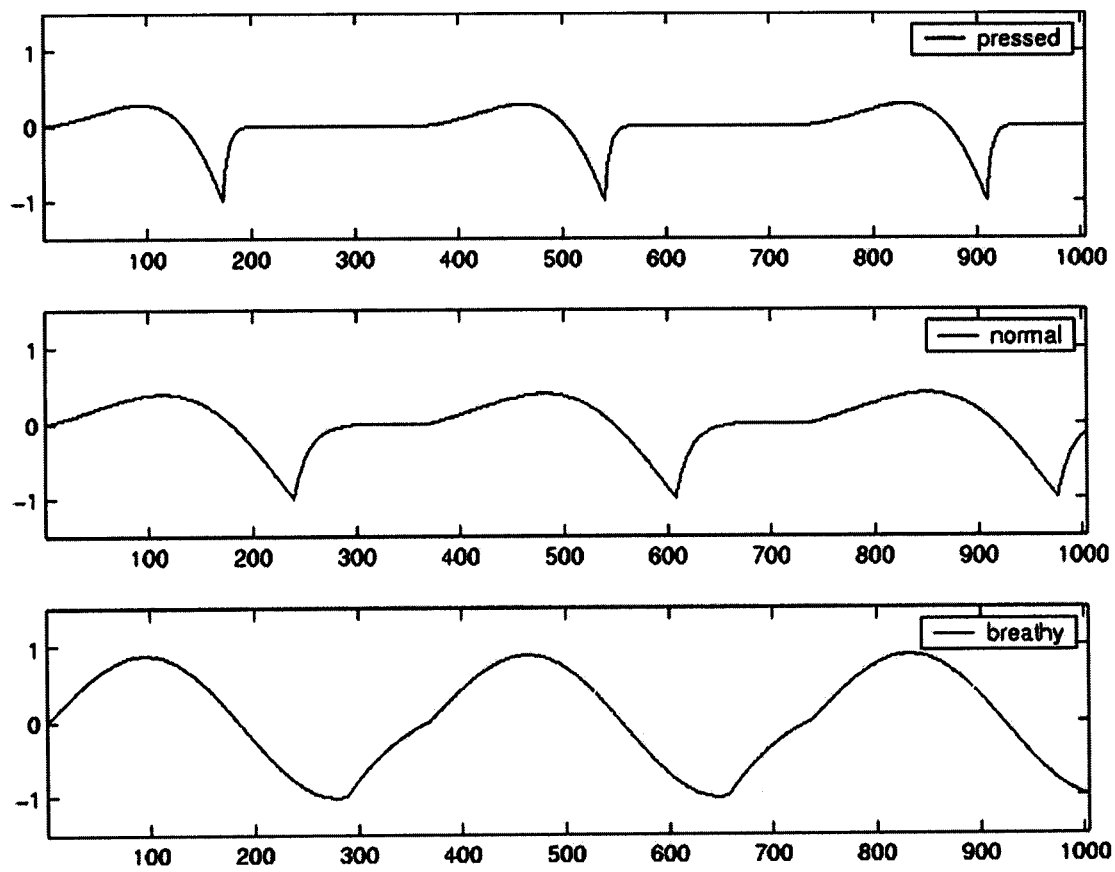
【Figure 6】



【Figure 7】

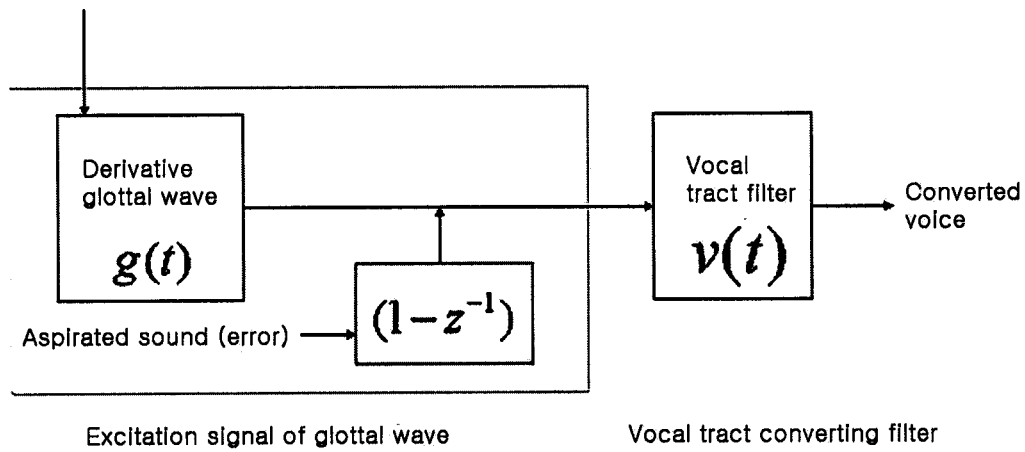


【Figure 8】

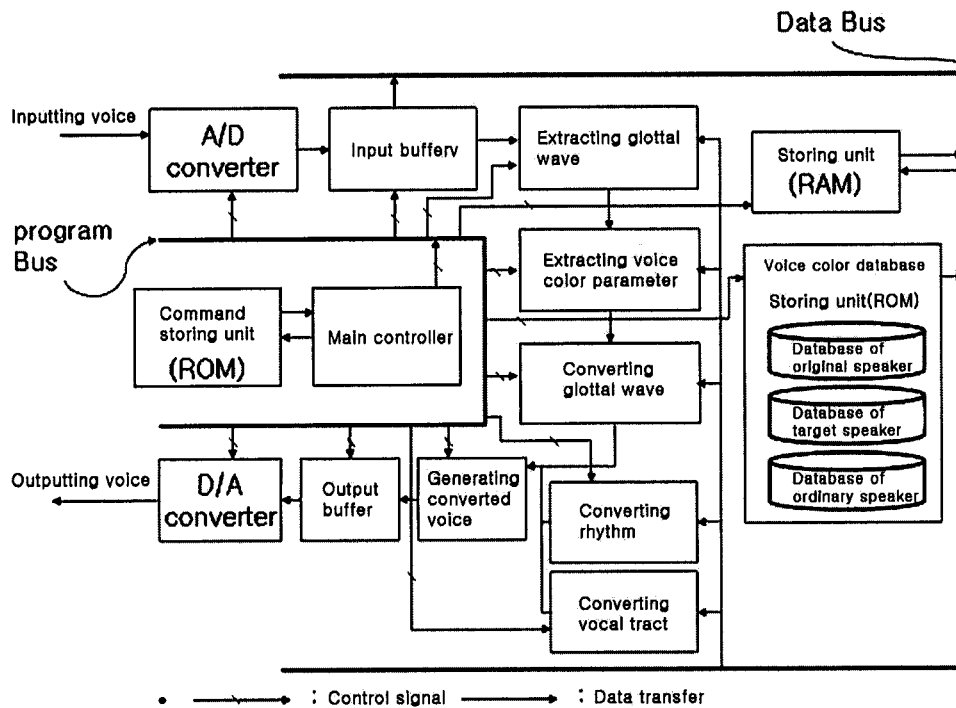


【Figure 9】

Basic frequency, Duration time, Energy



【Figure 10】



INTERNATIONAL SEARCH REPORT

International application No.
PCT/KR2006/004478**A. CLASSIFICATION OF SUBJECT MATTER***G10L 13/02(2006.01)i*

According to International Patent Classification (IPC) or to both national classification and IPC

B. FIELDS SEARCHED

Minimum documentation searched (classification system followed by classification symbols)

IPC 8 : G10L 15/00 ~ G10L 15/28, G10L 11/00 ~ G10L 13/08

Documentation searched other than minimum documentation to the extent that such documents are included in the fields searched

Utility models and applications for Utility Models since 1975

Japanese Utility Models and application for Utility Models since 1975

Electronic data base consulted during the international search (name of data base and, where practicable, search terms used)

eKIPASS(KIPO internal) "vocal, glot*, voice, color, conver*, trans*, source speaker, target speaker"

C. DOCUMENTS CONSIDERED TO BE RELEVANT

Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
A	US 6950799 B1 (NING BI et al.) 27 SEPTEMBER 2005 See the abstract, figures 1,5	1-12
A	US 6615174 B1 (ARSLAN L.M. et al.) 2 SEPTEMBER 2003 See the abstract, claims, figures 2-5	1-12
A	STYLIANOU, Y. et al. 'Continuous probabilistic transform for voice conversion' in IEEE Trans. on Speech and Audio Processing, Vol.6, No.2, March 1998. See the abstract	1-12
A	KAIN, A. et al. 'Spectral voice conversion for text-to-speech synthesis' in Proc. ICASSP, 1998, pp.285-288. See the abstract	1-12
A	KR 10-2004-61709 A (COREVOICE Co.) 7 JULY 2004 See the abstract, claims, figures 3-9	1-12



Further documents are listed in the continuation of Box C.



See patent family annex.

* Special categories of cited documents:

"A" document defining the general state of the art which is not considered to be of particular relevance

"E" earlier application or patent but published on or after the international filing date

"L" document which may throw doubts on priority claim(s) or which is cited to establish the publication date of citation or other special reason (as specified)

"O" document referring to an oral disclosure, use, exhibition or other means

"P" document published prior to the international filing date but later than the priority date claimed

"T" later document published after the international filing date or priority date and not in conflict with the application but cited to understand the principle or theory underlying the invention

"X" document of particular relevance; the claimed invention cannot be considered novel or cannot be considered to involve an inventive step when the document is taken alone

"Y" document of particular relevance; the claimed invention cannot be considered to involve an inventive step when the document is combined with one or more other such documents, such combination being obvious to a person skilled in the art

"&" document member of the same patent family

Date of the actual completion of the international search

23 APRIL 2007 (23.04.2007)

Date of mailing of the international search report

23 APRIL 2007 (23.04.2007)

Name and mailing address of the ISA/KR

Korean Intellectual Property Office
920 Dunsan-dong, Seo-gu, Daejeon 302-701,
Republic of Korea

Facsimile No. 82-42-472-7140

Authorized officer

KYUNG, Youn Jeong

Telephone No. 82-42-481-8536



INTERNATIONAL SEARCH REPORT

Information on patent family members

International application No.

PCT/KR2006/004478

Patent document cited in search report	Publication date	Patent family member(s)	Publication date
US06950799	27.09.2005	AU2003213179A1 CN1647159A MXPA04008005A US20030158728A1 W02003071523A1	09.09.2003 27.07.2005 26.11.2004 21.08.2003 28.08.2003
US06615174	02.09.2003	AT277405E AU6044298A1 DE69826446C0 DE69826446T2 EP970466A2 EP970466A4 EP970466B1 W09835340A2 W09835340A3	15.10.2004 26.08.1998 28.10.2004 20.01.2005 12.01.2000 31.05.2000 22.09.2004 13.08.1998 19.11.1998
KR 10-2004-61709 A	07.07.2004	None	