



(19) 대한민국특허청(KR)
(12) 등록특허공보(B1)

(45) 공고일자 2010년08월25일
(11) 등록번호 10-0978015
(24) 등록일자 2010년08월19일

(51) Int. Cl.
H04M 9/00 (2006.01) H04M 9/08 (2006.01)
G10L 21/02 (2006.01)
(21) 출원번호 10-2004-7021716
(22) 출원일자(국제출원일자) 2003년06월19일
심사청구일자 2008년06월19일
(85) 번역문제출일자 2004년12월31일
(65) 공개번호 10-2005-0016719
(43) 공개일자 2005년02월21일
(86) 국제출원번호 PCT/IB2003/002823
(87) 국제공개번호 WO 2004/004297
국제공개일자 2004년01월08일
(30) 우선권주장
02077608.4 2002년07월01일
유럽특허청(EPO)(EP)

(56) 선행기술조사문헌
W0199710586 A1
W0199745995 A1
US6108428 A
US6151397 A

전체 청구항 수 : 총 11 항

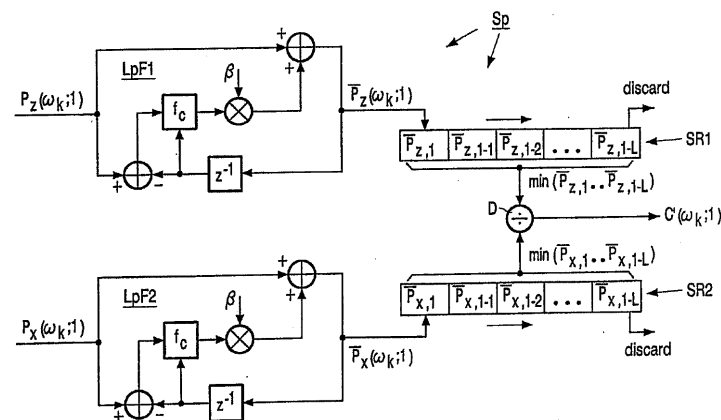
심사관 : 송병준

(54) 고정 스펙트럼 전력 의존 오디오 강화 시스템

(57) 요약

왜곡된 요구 신호(z)를 나르기 위한 신호 입력, 기준 신호 입력, 및 양쪽 신호 입력들에 결합되어 요구 신호의 왜곡에 대한 추정으로서 역할을 하는 기준 신호 x 에 의해 왜곡된 요구 신호를 처리하는 스펙트럼 프로세서(SP)를 포함하는 음성 인식이나 보이스 제어를 위한 오디오 강화 시스템(1)이 개시된다. 스펙트럼 프로세서(SP)는 인자 C' 가 결정되어, 추정이 기준 신호의 스펙트럼 전력에 대한 C' 곱의 함수가 되고, 인자 C' 는 시간에 대해 본질적으로 고정인 신호들(z, x)의 성분들 간의 스펙트럼 비율에 따라 결정되도록 상기 처리를 위해 제공된다. 상기 신호들의 고정 부분들에 의해 결정된 인수는 오디오 강화 시스템 내 중요 음성 검출기의 애플리케이션을 불필요하게 한다.

대표도



특허청구의 범위

청구항 1

왜곡된 요구 신호(distorted desired signal) z 를 위한 신호 입력,

기준 신호 x 를 위한 기준 신호 입력, 및

상기 두 신호 입력들 모두에 결합되고, 상기 왜곡된 요구 신호 z 의 왜곡 추정치를 도출되도록 하는 상기 기준 신호 x 에 의해 상기 왜곡된 요구 신호 z 를 처리하여, 실질적으로 왜곡 없는 출력 신호 q 를 산출하는 스펙트럼 프로세서를 포함하는 오디오 강화 시스템에 있어서,

상기 스펙트럼 프로세서는 인자 C' 을 이용하여 상기 왜곡된 요구 신호 z 를 처리를 위해 제공되고,

상기 왜곡된 요구 신호 z 의 상기 왜곡 추정치는 상기 인자 C' 과 상기 기준 신호 x 의 스펙트럼 전력(spectral power)의 곱인 함수이고,

상기 인자 C' 은 왜곡된 요구 신호 z 와 상기 기준 신호 x 의 본질적으로 고정인 성분들의 스펙트럼 비율에 의해 결정되는 것을 특징으로 하는, 오디오 강화 시스템.

청구항 2

제 1 항에 있어서,

상기 인자 C' 는 상기 왜곡된 요구 신호의 스펙트럼 전력 중 최소값과 상기 기준 신호의 스펙트럼 전력 중 최소값의 비에 의해 규정되고,

상기 두 최소값 모두는 특정 시간 기간 동안 결정되는 것을 특징으로 하는, 오디오 강화 시스템.

청구항 3

제 2 항에 있어서,

상기 시간 기간은 상기 왜곡된 요구 신호 내 적어도 하나의 휴지(pause)를 포함하는 것을 특징으로 하는, 오디오 강화 시스템.

청구항 4

제 3 항에 있어서,

상기 시간 기간은 적어도 4초 내지 5초를 지속하는 것을 특징으로 하는, 오디오 강화 시스템.

청구항 5

제 1 항 내지 제 4 항 중 어느 한 항에 있어서,

상기 스펙트럼 전력은 스펙트럼 진폭, 제곱 스펙트럼 크기, 전력 스펙트럼 밀도, 또는 멜-스케일 평활화된 스펙트럼 밀도(Mel-scale smoothed spectral density)와 같은, 상기 스펙트럼 전력에 관계된 어떤 양(陽)의 함수에 의해 규정되는 것을 특징으로 하는, 오디오 강화 시스템.

청구항 6

제 1 항 내지 제 4 항 중 어느 한 항에 있어서,

상기 스펙트럼 프로세서는 스펙트럼 전력의 값들을 저장하는 쉬프트 레지스터들을 포함하는 것을 특징으로 하는, 오디오 강화 시스템.

청구항 7

제 1 항 내지 제 4 항 중 어느 한 항에 있어서,

상기 스펙트럼 전력은 평활화된 스펙트럼 전력인 것을 특징으로 하는, 오디오 강화 시스템.

청구항 8

제1항에 따른 오디오 강화 시스템이 제공된 통신 시스템.

청구항 9

왜곡된 요구 신호 z 를 강화하는 방법으로서,

상기 왜곡된 요구 신호 z 의 왜곡 추정치를 도출되도록 하는 기준 신호 x 에 의해 상기 왜곡된 요구 신호 z 를 처리하여, 실질적으로 왜곡 없는 출력 신호 q 를 산출하는 처리 단계를 포함하는 왜곡된 요구 신호 z 를 강화하는 방법에 있어서,

상기 처리 단계는 인자 C' 을 이용하여 상기 왜곡된 요구 신호 z 를 처리하는 단계를 포함하고,

상기 왜곡된 요구 신호 z 의 상기 왜곡 추정치는 상기 인자 C' 과 상기 기준 신호 x 의 스펙트럼 전력(spectral power)의 곱인 함수이고,

상기 인자 C' 은 왜곡된 요구 신호 z 와 상기 기준 신호 x 의 본질적으로 고정인 성분들의 스펙트럼 비율에 의해 결정되는 것을 특징으로 하는, 왜곡된 요구 신호 강화 방법.

청구항 10

삭제

청구항 11

제1항에 따른 오디오 강화 시스템이 제공된 음성 인식 시스템.

청구항 12

제1항에 따른 오디오 강화 시스템이 제공된 보이스 제어 시스템.

명세서

기술분야

[0001] 본 발명은 왜곡된 요구 신호(distorted desired signal)를 나르기 위한 신호 입력, 기준 신호를 위한 기준 신호 입력(reference signal input), 및 상기 두 신호 입력들에 결합되어 요구 신호의 왜곡에 대한 추정으로서 역할을 하는 기준 신호에 의해 상기 왜곡된 요구 신호를 처리하는 스펙트럼 프로세서(spectral processor)를 포함하는 오디오 강화 시스템에 관한 것이고, 거기에 사용하는데 적합한 신호들에 관한 것이다.

[0002] 본 발명은 이동 전화, 음성 인식 시스템, 또는 보이스 제어 시스템 등의 핸드프리 통신 장치, 특히 통신 시스템과 같은 시스템에 관한 것이고, 상기 시스템은 오디오 강화 시스템이 제공된다. 왜곡된 요구 신호를 강화하는 방법으로서, 상기 신호가 상기 요구 신호의 왜곡에 대한 추정으로서 역할을 하는 기준 신호에 의해 스펙트럼 처리되는, 상기 방법에 관한 것이다.

배경기술

[0003] 상기 오디오 강화 시스템은 W0 97/45995로부터 알려진 왜곡 잡음과 같은 간섭 성분을 억제하기 위한 장치에 의해 구현된다. 알려진 시스템은 오디오 신호 입력들에 결합된 다수의 마이크로폰들을 포함한다. 마이크로폰들은 왜곡된 요구 신호를 위한 주요 마이크로폰과, 간섭 신호를 수신하기 위한 하나 이상의 참조 마이크로폰들을 포함한다. 상기 시스템은 또한, 오디오 신호 입력들을 통해 마이크로폰들에 결합된 신호 처리 장치로 구현된 스펙트럼 프로세서를 포함한다. 시스템의 출력에 감소된 간섭 잡음 성분을 포함한 출력 신호를 나타내도록, 왜곡 신호로부터 간섭 신호를 스펙트럼 감산한다.

[0004] 알려진 오디오 강화 시스템의 단점은, 그것의 간섭 신호 소거 능력들이 음성 처리 장치에 결합될 음성 검출기의 애플리케이션에 의존한다는 점이다. 알려진 오디오 강화 시스템의 동작은 결정적으로 상기 음성 검출기에 의한 적절한 음성 검출에 의존한다.

발명의 상세한 설명

- [0005] 그러므로, 본 발명의 목적은 개선된 오디오 강화 시스템과, 음성 검출기의 존재와 중요 동작에 의해 복잡해지지 않는 관련 방법을 제공하는 것이다.
- [0006] 또한, 본 발명에 따른 오디오 강화 시스템은, 스펙트럼 프로세스가 인자 C' 가 결정되어 추정이 기준 신호의 스펙트럼 전력에 대한 C' 곱이고, 인자 C' 가 본질적으로 시간에 대해 고정인 신호 z 및 x 의 성분들 간의 스펙트럼 비율에 따라 결정되도록 상기 처리를 위해 제공되는 것을 특징으로 한다.
- [0007] 유사하게, 본 발명에 따른 방법은, 상기 추정이 기준 신호의 스펙트럼 전력에 대한 C' 곱의 함수가 되고, C' 가 시간에 대해 본질적으로 고정인 신호들 z 및 x 의 성분들 간의 스펙트럼 비율에 따라 결정되는 것을 특징으로 한다.
- [0008] 발명자는, 규정된 바와 같은 인자 C' 가 요구 신호에 본질적으로 둔감하다는 것을 발견했다. 인자 C' 만이 신호들 z 및 x 내 고정 성분들의 비를 차지한다. 상기 인자 C' 의 개념으로, 음성 검출기를 사용할 필요 없이 오디오 강화 시스템에 실제로 입력되는 왜곡된 요구 신호의 왜곡에 대해 신뢰성 있는 추정이 제공될 수 있다. 이것은 본 발명에 따른 단순화된 오디오 강화 시스템의 개선되고 덜 중요한 왜곡 소거 특성들을 유발한다. 하나 이상의 기준 신호들이 잡음, 에코들, 경합 음성, 요구 음성 신호의 잔향 등과 같은 왜곡들을 포함하는 경우, 왜곡 소거가 특히 개선된다. 또한, 왜곡에 대한 주파수 의존 추정은 일부 기준 신호(들)이 이용가능한 임의의 시나리오에서 계산될 수 있다.
- [0009] 다른 이점들은 잡음 플로어(noise floor) 또는 에코 테일(echo tail)과 같은 개개의 왜곡 성분들의 명백한 추정이 필요치 않다는 점이고, 요구된다면, 이들 성분들을 갖는 조합 기술이 쉽게 달성될 수 있다. 이것은 마이크로폰 빔 형성 애플리케이션들을 위해서와 같이 우수한 추정 기술들이 존재하지 않는 경우에 특히 바람직하다.
- [0010] 본 발명에 따른 오디오 강화 시스템의 실시예는 제 2 항에 약속된 독특한 특징들을 갖는다.
- [0011] 일반적으로, 특정한 수의 시간 프레임들을 커버하는 평균 스펙트럼 전력의 형태를 갖는 스펙트럼 전력들이 측정된다. 일정 시간 기간 동안, 본 발명에 따른 오디오 강화 시스템의 계산 복잡성에 대한 실제 부담 없이, 두 스펙트럼 전력들 모두에 대한 시간 기간 최소값들이 결정된다.
- [0012] 본 발명에 따른 오디오 강화 시스템의 다른 실시예에서, 시간 기간은 왜곡된 요구 신호 내 적어도 하나의 휴지(pause)를 포함한다. 이것은 잘 결정된 최소값과 왜곡된 요구 입력 신호의 고정 스펙트럼 성분 값을 유발하는데, 최소값은 입력 신호에서의 고정 왜곡을 정확히 표현한다.
- [0013] 바람직하게, 오디오 강화 시스템에 대한 왜곡된 요구 신호 입력에서의 음성 휴지가 정상적으로 포함되도록 상기 시간 기간은 4초 내지 5초를 지속한다.
- [0014] 오디오 강화 시스템의 또 다른 실시예가 제 5 항에 약속된 독특한 특징들을 갖는다.
- [0015] 일반적으로, 요구 신호의 왜곡 추정은 예를 들어, 신호 전력 또는 신호 에너지에 의해 양의 함수로 바람직하게 표현될 수 있고, 상기 스펙트럼 유닛들 중 하나에 의해 규정된다.
- [0016] 실질적으로 양호한 실시예는 제 6 항의 독특한 특징들을 갖는다.
- [0017] 상기의 경우, 오디오 강화 시스템은 스펙트럼 전력들 및/또는 평활화된 스펙트럼 전력들의 값들을 저장하기 위한 쉬프트 레지스터들을 구현하기 위해 비용 효율적이고 용이하게 포함된다.
- [0018] 이제, 본 발명에 따른 오디오 강화 시스템과 방법은 첨부된 도면을 참조하면서, 부가적인 이점들과 함께 더 상세히 설명될 것이며, 여기서 유사한 성분들은 동일한 참조 번호들에 의해 언급된다.

실시예

- [0022] 도 1은 스펙트럼 프로세서 SP에 의해 구현된 오디오 강화 시스템(1)의 기본 다이어그램을 도시하며, 여기에 주파수 도메인 입력 신호들(z , x) 및 출력 신호(q)가 도시된다. 이들 주파수 도메인 신호들은, 간단히 STFT로 언급되는 단시간(short time) DFT와 같은 이산 푸리에 변환(Discrete Fourier Transform)에 의해 상기 프로세서 SP에서 블록 단위로(block-wise) 스펙트럼 계산된다. 이 STFT는 인수들 kB , lw_0 나 때때로 인수 w_k 만으로 표현될 수 있는 시간 및 주파수 모두의 함수이다. 여기서 k 는 이산 시간 프레임 인덱스를 나타내고, B 는 프레임 쉬프트를 나타내며, l 은 (이산) 주파수 인덱스를 나타내고, w_0 는 기본 주파수 스페이싱을 나타내며, w_k 는 인덱스 k 의 스펙트럼 성분을 나타낸다. 입력 신호 z 는 왜곡된 요구 신호를 나타낸다. 그것은 일반적으로 음성의 형태

로, 요구 신호의 잡음, 에코, 경합 음성 또는 잔향과 같은 왜곡들과 요구 신호와의 합을 포함한다. 신호 x 는 왜곡된 요구 신호 z 내의 왜곡 추정치 도출되는 기준 신호를 나타낸다. 신호들(z, x)은 도 1과 도 2에 도시된 바와 같이, 하나 이상의 마이크로폰들(2)로부터 발생할 수 있다. 다중-마이크로폰 오디오 강화 시스템(1)에서, 하나 이상의 마이크로폰들로부터 기준 신호를 도출하기 위해 두 개 이상의 별개의 마이크로폰들(2)이 존재한다.

[0023] 오디오 강화 시스템(1)은 그것으로부터 기준 신호 x 를 도출하기 위해 적응 필터 수단(도시되지 않음)을 포함할 수 있다. 이 경우, 기준 신호는 통신 시스템의 최단부에서 발생한다.

[0024] 도 1의 실시예에서, 신호 x 만이 기준 또는 잡음 신호를 포함하는 반면, 신호 z 는 요구 신호 및 잡음 신호 모두를 포함한다. 도 2는 마이크로폰(2)들 모두가 마이크로폰 어레이 신호들(u_1, u_2)을 통해 음성과 잡음을 감지하는 경우에 대한 오디오 강화 시스템(1)의 실시예를 도시한다. 필터와 합계 빔 형성기(sum beamformer)(3)는 이제, 마이크로폰들(2) 및 스펙트럼 프로세서(SP) 사이에 결합된다. 또한, 스펙트럼 프로세서는 전송된 신호들(z, x)을 수신하는데, 신호(x)만이 기준 또는 잡음을 포함하고, 신호(z)는 요구 신호 및 잡음 신호를 포함한다. 개개의 전달 함수들 $f_1(w), f_2(w)$ 을 통해, 왜곡된 요구 신호(z)가 마이크로폰 어레이 신호들(u_1, u_2) 각각의 선형 조합에 의해 획득되도록 상기 빔 형성기(3)가 설계된다. 요구 신호에 직교(orthogonal)하는 서브스페이스로 이들 신호들을 투사하기 위해, 기준 신호 x 는 블로킹 행렬(blocking matrix) $B(w)$ 에 의해 개개의 마이크로폰 어레이 신호들로부터 도출된다. 이상적으로, 행렬 $B(w)$ 의 출력 신호(x)는 요구 음성을 포함하지 않지만 왜곡들을 포함한다. 다음, 기준 신호(x)에 의해 상기 왜곡된 요구 신호(z)를 스펙트럼 처리하기 위해 신호들(z, x)이 스펙트럼 프로세서(SP)에 공급된다. 상기 프로세서(SP)로부터의 신호(q)는 사실상 왜곡 없는 출력 신호이다. $q=G \times z$ 인데, 여기서 G 는 이하에서 기술될 이득 함수이다.

[0025] 오디오 강화 시스템(1)은 시스템, 특히 이동 전화, 음성 인식 시스템, 또는 보이스 제어 시스템 등의 핸드프리 통신 장치와 같은 통신 시스템에 포함될 수 있다.

[0026] 스펙트럼 프로세서(SP)는 전송된 이산 푸리에 변환(DFT)에 의해 발생된 후속 주파수 빈들에 대한 제어가능한 이득 함수로서 역할을 하도록 동작한다. 이러한 이득 함수는 왜곡된 요구 음성 신호(z)에 적용되지만, 신호(z)의 위상은 변경되지 않는다. 우수한 성능의 오디오 강화를 위해, 이득 함수형, 즉, 입력 신호에 존재하는 왜곡의 추정은 중요하다. 그러나, 다양한 이득 함수들이 처리되는 최적화 기준에 종속하여 사용될 수 있다. 예들은 스펙트럼 감산(spectral subtraction), 위너 필터링(Wiener filtering) 또는, 예를 들어 스펙트럼 진폭이나 크기에 기초한 MMSE 추정(Minimum Mean-Square Error estimation)이나 로그-MMSE 추정(log-MMSE estimation), 제곱 스펙트럼 크기(squared spectral magnitude), 전력 스펙트럼 밀도(power spectral density), 또는 포함된 신호들의 멜-스케일 평활화된 스펙트럼 밀도(Mel-scale smoothed spectral density)를 포함한다. 이들 기술들은 하나 이상의 마이크로폰들 및/또는 라우드스피커들을 갖는 오디오 강화 시스템들(1)에 대해 전송된 애플리케이션들과 결합될 수 있다.

[0027] 예를 들어, 이하에 기술될 위너 필터형의 경우, 스펙트럼 프로세서(SP)에서 구현된 이득 함수는 이하의 형태를 갖는다:

$$G(kB, lw_0) = 1 - \gamma P_{zz,n}(kB, lw_0)/P_{zz}(kB, lw_0) \quad (1)$$

[0029] 여기서 $P_{zz,n}(kB, lw_0)$ 와 $P_{zz}(kB, lw_0)$ 는 입력 신호(z)에서의 왜곡의 전력 분배와, 입력 신호(z) 자체의 전력 분배에 대한 추정치들이다. γ 는 왜곡에 적용된 억제량을 조정하게 하는 소위 초과 감산 인자를 나타낸다. 이처럼 트레이드-오프(trade-off)는 왜곡 억제량과 프로세서 출력 신호의 지각 품질 사이에 형성될 수 있다.

[0030] 식(1)에서, $P_{zz,n}(kB, lw_0)$ 는 일반적으로 알려지지 않았고, 따라서 추정되어야 한다. 추정 $\hat{P}$ 이 이하와 같이 제안된다:

$$P_{zz,n}(kB, lw_0) = C(kB, lw_0) * P_{xx}(kB, lw_0) \quad (2)$$

[0032] 여기서 비율 항은:

$$C(kB, lw_0) = P_{zz}(kB, lw_0)/P_{xx}(kB, lw_0). \quad (3)$$

[0034] 여기서 $P_{zz}(kB, lw_0)$ 은 음성과 같은 요구 신호의 부재 동안 측정된, 왜곡된 요구 신호(z)에 대한 왜곡의 시평균

스펙트럼 전력(time averaged spectral power)이고, $P_{xx}(kB, lw_0)$ 은 기준 신호(x)의 시평균 스펙트럼 전력이다. 예를 들면, 스펙트럼 진폭 또는 크기와 같이 스펙트럼 전력에 대한 양(陽, positive)의 측정값으로서, 포함된 신호들의 제곱 스펙트럼 크기, 전력 스펙트럼 밀도, 또는 멜-스케일 평활화된 스펙트럼 밀도가 취해질 수 있다. 프로세서(SP)에서 식(3)의 구현은 음성 검출기를 요구한다. 상기 음성 검출기가 정확히 수행하지 않는다면, 요구 음성이 영향받을 수 있고, 이것은 방지되어야 하는 가청 아티팩트들(audible artifacts)을 초래한다. 그러나, 차나 공장과 같은 잡음 조건들에서 신뢰성 있는 음성 검출이 태스크를 수행하는 것은 어렵다.

[0035] 일반적으로, 인자 C에 대한 추정으로서, 실제로 식(3)의 비 중 고정 부분들에 집중함으로써 행해지는, 요구 음성에 사실상 둔감한 새로운 인자 C'를 제안함으로써, 음성 검출기를 요구하지 않는 강력한 알고리즘이 생성된다. 상기 아이디어의 실제적 구현에 있어서, 인자 C'는 왜곡된 요구 신호(z)의 최소 스펙트럼 전력과 기준 신호(x)의 최소 스펙트럼 전력의 비에 의해 규정되는데, 두 최소값은 하나의 시간 기간 동안 결정된다. 이하와 같은 공식으로 표현된다:

[0036]
$$C'(w_k; l) = \min_{m \in [1-L, \dots, l]} P_{zz}(w_k; m) / \min_{m \in [1-L, \dots, l]} P_{xx}(w_k; m). \quad (4)$$

[0037] 1-L개의 시간 프레임들과 1개의 시간 프레임들 간의 시간 기간은 왜곡된 요구 신호에 존재하는 적어도 하나의 휴지(pause)를 포함하는 L개의 시간 프레임들을 커버한다. 요구 신호가 음성 신호라면, 일반적으로 이것은 음성 휴지이다. 그렇게 결정된 최소값은 신호들(z, x) 각각의 고정 성분들에 대해 식 (4)의 비를 집중시키며, 이때 최소값은 왜곡 또는 잡음의 고정 성분들을 나타낸다. 일반적으로 시간 기간은 적어도 4 내지 5초를 지속한다. 식(4)에 주어진 인자 C'는 신호들(z, x)의 고정 성분들에 기초하여 결정된다. 음성과 같은 비 고정 성분들이 이 신호들에 존재하는 경우에도 지속되는 것이 가정되고, 스펙트럼 프로세서(SP)에 의해 수행된 동작은 상기 가정에 기초한다.

[0038] 식(4)에서 인자 C'의 분자와 분모의 스펙트럼은 평활화 상수들(β)을 각각 갖는 블록들(LPF1, LPF2)에 구현된 1차 재귀들(recursions)에서 전력 스펙트럼을 평활화함으로써 획득된다. 이 블록들에서의 재귀 구현은 입력 x, z 신호들의 평활화된 전력 스펙트럼 밀도 버전들을 획득하도록 도식된 바와 같이 결합된 곱셈기들(X), 덧셈기들(+), 및 지연 선들(z^{-1})을 포함한다. 예를 들어, z 신호는 이 후, 평활화 규칙을 준수한다:

[0039]
$$P_{zz}(w_k; l) = \beta P_{zz}(w_k; l) + (1-\beta) P_{zz}(w_k; l-1)$$

[0040] 여기서 평활화 상수(β)는 0과 1 사이의 값으로 가정한다. 동일한 규칙이 x 신호 스펙트럼에 대해 적용할 수 있다. β 의 값은 임의의 요구 방식으로 제어될 수 있다. 그것의 값은 일반적으로 50-200ms의 시간 상수에 대응한다. 시간 프레임 인덱스 m마다, 이들 평활화된 양들의 각각이 쉬프트 레지스터들 SR1, SR2의 형태로 각각 버퍼에 저장된다. 레지스터 위치들의 각각에 저장된 L개의 평활화된 값들 외에, 개개의 최소 값들이 식(4)에 따른 C'의 계산된 값을 나타내도록 제수(divisor) D에 공급된다. 물론, 분모에서 작은 값에 의한 나눗셈을 방지하도록 적절한 조치들이 취해진다.

[0041] 요구 음성 신호의 평균 레벨이 왜곡의 평균 레벨에 대해 너무 높은 경우, LPF1과 LPF2에 의한 평균 출력이 요구 음성에 의해 지배되는 문제가 발생할 수 있다. 이것은 이들 평균들이 고 음성 레벨의 발생 후에 저 왜곡 레벨로 복귀하는데 오랜 시간이 걸리기 때문이다. C'의 추정이 요구 음성에 의해 영향을 받을 수 있을 때, 요구 음성 신호의 원하지 않는 억제를 초래한다. 이러한 영향은 예를 들어 이하에 따라 재귀들 내 다중가변 압축 함수(multivariable compression function) f_c 를 적용함으로써 감소된다:

[0042]
$$P_{zz}(w_k; l) = \beta P_{zz}(w_k; l-1) + (1-\beta) f_c\{P_{zz}(w_k; l), P_{zz}(w_k; l-1)\}$$

[0043] 동일한 규칙이 x 신호에 대해 적용할 수 있다. 새로운 입력 전력 값이 필터들(LPF1, LPF2)에서의 값들에 비해 비교적 큰 경우, 재귀의 업데이트 단계를 줄이기 위해 압축 함수가 선택된다. 그러므로 압축 함수는 평균 신호 전력상에서 높은 요구 음성 레벨의 영향을 감소시킨다. 적합한 압축 함수의 예가 이하에서 주어진다:

[0044]
$$f_c(A, B) = \min \{A-B, \delta B\}$$

[0045] 여기서 δ 는 양의 상수이다. δ 의 값이 작을수록, 신호 값에서 더 느린 상승이 재귀 필터들(LPF1, LPF2)에 이어진다. 압축 블록(f_c)을 포함하는 실시예가 도 3에 도시된다. 압축이 필요치 않다면, 간단히 생략될 수 있다.

[0046] 상기는 본질적으로 양호한 실시예들과 최상의 가능 모드들을 참조하여 기술되었지만, 첨부된 청구범위 내의 다양한 변경들, 특징들, 및 특징들의 조합이 당업자의 이해 범위 내에 있기 때문에, 상기 실시예들은 관련된 시스템 실시예들 및 방법의 예들을 제한하는 것으로 해석되지 않는다는 점이 이해될 것이다.

도면의 간단한 설명

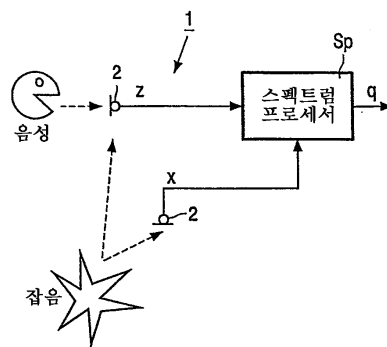
[0019] 도 1은 본 발명에 따른 오디오 강화 시스템의 기본 다이어그램.

[0020] 도 2는 필터 및 합계 빔 형성기를 갖는 본 발명에 따른 오디오 강화 시스템의 또 다른 실시예에 구현된 기본 다이어그램.

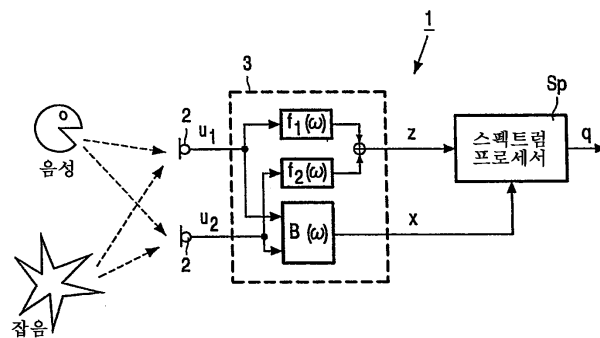
[0021] 도 3은 본 발명에 따른 오디오 강화 시스템의 상세한 실시예를 도시하는 도면.

도면

도면1



도면2



도면3

