



(12)发明专利

(10)授权公告号 CN 103890843 B

(45)授权公告日 2017.01.18

(21)申请号 201280051123.7

(22)申请日 2012.10.16

(65)同一申请的已公布的文献号

申请公布号 CN 103890843 A

(43)申请公布日 2014.06.25

(30)优先权数据

61/548998 2011.10.19 US

(85)PCT国际申请进入国家阶段日

2014.04.17

(86)PCT国际申请的申请数据

PCT/IB2012/055628 2012.10.16

(87)PCT国际申请的公布数据

W02013/057659 EN 2013.04.25

(73)专利权人 皇家飞利浦有限公司

地址 荷兰艾恩德霍芬

(72)发明人 P.科奇奇安 S.斯里尼瓦桑

(74)专利代理机构 中国专利代理(香港)有限公司 72001

代理人 孙之刚 汪扬

(51)Int.Cl.

G10L 21/0216(2013.01)

(56)对比文件

US 2004/0167777 A1,2004.08.26,

CN 1530928 A,2004.09.22,

US 2008/0140396 A1,2008.06.12,

US 7478043 B1,2009.01.13,

Sriram Srinivasan et al..Codebook

Driven Short-Term Predictor Parameter

Estimation for Speech Enhancement.《IEEE

TRANSACTIONS ON AUDIO, SPEECH, AND

LANGUAGE PROCESSING》.2006,第14卷(第1期),

163-176.

审查员 康丹丹

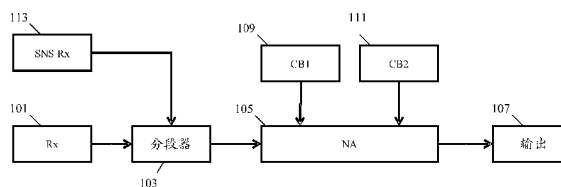
权利要求书2页 说明书14页 附图3页

(54)发明名称

信号噪声衰减

(57)摘要

一种噪声衰减装置,其接收包括预期和噪声信号分量的第一信号。两个码本(109,111)分别包括分别表示可能的预期和噪声信号分量的预期信号候选和噪声信号候选。噪声衰减器(105)通过为预期和噪声信号候选的每个配对生成作为预期信号候选与噪声信号候选的组合的估计信号候选来生成估计信号候选。然后从估计信号候选确定信号候选并且第一信号基于该信号候选被噪声补偿。使用表示对环境中的预期来源或噪声的测量的传感器信号来降低所搜索的候选数量,从而大幅降低复杂度以及运算资源使用。特别地,噪声衰减可以是音频噪声衰减。



1. 一种噪声衰减装置,包括:

-接收机(101),其用于接收针对环境的第一信号,所述第一信号包括对应于来自环境中的预期来源的信号的预期信号分量,以及对应于环境中的噪声的噪声信号分量;

-第一码本(109),其包括用于预期信号分量的多个预期信号候选,每一个预期信号候选表示可能的预期信号分量;

-第二码本(111),其包括用于噪声信号分量的多个噪声信号候选,每一个噪声信号候选表示可能的噪声信号分量;

-输入(113),其用于接收提供对环境的测量的传感器信号,所述传感器信号表示对环境中的预期来源或噪声的测量;

-分段器(103),其用于将第一信号分段到时间段中;

-噪声衰减器(105),其包括被设置成为每个时间段执行以下步骤:

-通过为第一码本的码本条目的第一群组中的预期信号候选与第二码本的码本条目的第二群组中的噪声信号候选的每个配对生成组合信号来生成多个估计信号候选;

-从估计信号候选中生成用于在该时间段中的第一信号的信号候选,以及

-响应于信号候选来衰减该时间段中的第一信号的噪声;

其中噪声衰减器(105)被设置成通过响应于参考信号来选择码本条目的子集而生成第一群组和第二群组中的至少一个。

2. 权利要求1的噪声衰减装置,其中传感器信号表示对预期来源的测量,并且噪声衰减器(105)被设置成通过从第一码本(109)中选择码本条目的子集来生成第一群组。

3. 权利要求2的噪声衰减装置,其中第一信号是音频信号,预期来源是音频源,预期信号分量是语音信号,并且传感器信号是骨传导麦克风信号。

4. 权利要求2的噪声衰减装置,其中与预期信号分量相比,传感器信号提供精度较低的预期来源表示。

5. 权利要求1的噪声衰减装置,其中传感器信号表示对噪声的测量,并且噪声衰减器(105)被设置成通过从第二码本(111)中选择码本条目的子集来生成第二群组。

6. 权利要求5的噪声衰减装置,其中传感器信号是机械振动检测信号。

7. 权利要求5的噪声衰减装置,其中传感器信号是加速度计信号。

8. 权利要求1的噪声衰减装置,还包括映射器(301),其用于生成多个传感器信号候选与第一码本和第二码本中的至少一个的码本条目之间的映射;并且其中噪声衰减器(105)被设置成响应于所述映射来选择码本条目的子集。

9. 权利要求8的噪声衰减装置,其中噪声衰减器(105)被设置成响应于多个传感器信号候选中的每一个与传感器信号之间的距离量度来从多个传感器信号候选中选择第一传感器信号候选,并且响应于用于第一信号候选的映射来生成所述子集。

10. 权利要求8的噪声衰减装置,其中映射器(301)被设置成基于来自发出第一信号的输入传感器和发出传感器信号的传感器的同时测量来生成所述映射。

11. 权利要求8的噪声衰减装置,其中映射器(301)被设置成基于传感器信号候选与第一码本和第二码本中的至少一个的码本条目之间的差量度来生成所述映射。

12. 权利要求1的噪声衰减装置,其中第一信号是来自第一麦克风的麦克风信号,并且传感器信号是来自远离第一麦克风的第二麦克风的麦克风信号。

13. 权利要求1的噪声衰减装置, 其中第一信号是音频信号, 并且传感器信号来自非音频传感器。

14. 一种噪声衰减的方法, 包括:

- 接收针对环境的第一信号, 所述第一信号包括对应于来自环境中的预期来源的信号
的预期信号分量以及对应于环境中的噪声的噪声信号分量;

- 提供包括用于预期信号分量的多个预期信号候选的第一码本(109), 每一个预期信号
候选表示可能的预期信号分量;

- 提供包括用于噪声信号分量的多个噪声信号候选的第二码本(111), 每一个噪声信号
候选表示可能的噪声信号分量;

- 接收提供对环境的测量的传感器信号, 所述传感器信号表示对环境中的预期来源或
噪声的测量;

- 将第一信号分段到时间段中;

- 为每一个时间段, 执行以下步骤:

- 通过为第一码本的码本条目的第一群组中的预期信号候选与第二码本的码本条目的
第二群组中的噪声信号候选的每个配对生成组合信号来生成多个估计信号候选,

- 从估计信号候选中生成用于在该时间段中的第一信号的信号候选, 以及

- 响应于信号候选来衰减该时间段中的第一信号的噪声;

以及通过响应于参考信号来选择码本条目的子集而生成第一群组和第二群组中的至
少一个。

信号噪声衰减

技术领域

[0001] 本发明涉及信号噪声衰减,并且特别但不排他地涉及用于音频特别是语音信号的噪声衰减。

背景技术

[0002] 为了进一步增强或突出预期信号分量,在很多应用中信号中的噪声的衰减是所期望的。特别地,在很多场景中,音频噪声的衰减都是所期望的。例如,由于具有实际相关性,在存在背景噪声的情况下的语音增强业已引起了极大的关注。

[0003] 一种音频噪声衰减的方法是将两个或更多个麦克风的阵列与适当的波束形成算法一起使用。然而,这样的算法并不总是可行,或者提供了次优的性能。例如,它们往往对资源要求高,并且需要复杂的算法以用于追踪预期声源。此外,它们往往提供次优的噪声衰减,特别是在回响的和扩散非稳态噪声场中或者在存在许多干扰源的情况下更是如此。在这样的场景中,诸如波束形成之类的空间滤波技术只能取得有限的成功,并且在后处理步骤中通常会对波束形成器的输出执行附加的噪声抑制。

[0004] 业已提出的各种噪声衰减算法包含基于关于预期信号分量和噪声信号分量特性的认识或假设的系统。特别地,在非稳态噪声条件下,即便操作在单个麦克风信号上,诸如码本驱动方案之类的基于认识的语音增强方法也已显示出其执行良好。在以下文献中呈现了这样的方法的示例:S. Srinivasan、J. Samuelsson和W. B. Kleijn在2006年1月发表于IEEE Trans. Speech, Audio and Language Processing(IEEE会报:语音、音频和语言处理)第14卷第1号第163-176页的“Codebook driven short-term predictor parameter estimation for speech enhancement(用于语音增强的码本驱动短程预测器参数估计)”,以及S. Srinivasan、J. Samuelsson和W. B. Kleijn在2007年2月发表于IEEE Trans. Speech Audio Processing(IEEE会报:语音音频处理)第15卷第2号第441-452页的“Codebook based Bayesian speech enhancement for non-stationary environments(用于非稳态环境的基于码本的贝叶斯语音增强)”。

[0005] 这些方法依赖于经训练的语音和噪声谱形的码本,所述噪声谱形由例如线性预测(LP)系数来参数化。语音码本的使用是直观的并且容易有助于实际实现。语音码本可以要么与说话者无关(使用来自若干个说话者的数据来进行训练),要么与说话者相关。对于例如移动电话应用来说,由于这些应用往往是私人的,并且通常主要由单个说话者使用,因此后一种情况是有用的。然而,由于在实践中可能遇到各种噪声类型,因此,噪声码本在实际实现中的使用是具有挑战性的。结果,典型地使用非常大的噪声码本。

[0006] 典型地,基于这样的码本的算法试图找到当被结合时与所捕获的信号最接近匹配的语音码本条目和噪声码本条目。在已经找到恰当的码本条目时,该算法基于码本条目来补偿所接收的信号。然而,为了标识恰当的码本条目,在语音码本条目与噪声码本条目的所有可能组合之上执行搜索。这导致在运算上对资源要求非常高的过程,而这通常是不切实际的,对于低复杂度的设备来说尤其如此。另外,可能信号特别是噪声候选的大数量可能增

加错误估计的风险,从而导致次优的噪声衰减。

[0007] 因此,改进的噪声衰减方法将是有利的,并且特别地,允许增加的灵活性、降低的运算需求、便利的实现和/或操作、降低的成本和/或改善的性能的方法将是有利的。

发明内容

[0008] 因此,本发明试图优选地单个或以任何组合地缓解、减轻或消除以上所提及的一个或多个缺陷。

[0009] 根据本发明的一个方面,提供了一种噪声衰减装置,其包括:接收机,其用于接收针对环境的第一信号,所述第一信号包括对应于来自环境中的预期来源的信号预期信号分量,以及对应于环境中的噪声的噪声信号分量;第一码本,其包括用于预期信号分量的多个预期信号候选,每一个预期信号候选都表示可能的预期信号分量;第二码本,其包括用于噪声信号分量的多个噪声信号候选,每一个预期信号候选都表示可能的噪声信号分量;输入,其用于接收提供对环境的测量的传感器信号,该传感器信号表示对环境中的预期来源或噪声的测量;分段器,其用于将第一信号分段到时间段中;噪声衰减器,其包括被设置成为每个时间段执行以下步骤:通过为第一码本的码本条目的第一群组中的预期信号候选与第二码本的码本条目的第二群组中的噪声信号候选的每个配对生成组合信号来生成多个估计信号候选;从估计信号候选中生成用于该时间段中的第一信号的信号候选,以及响应于该信号候选衰减在该时间段中第一信号的噪声;其中噪声衰减器被设置成通过响应于参考信号来选择码本条目的子集而生成第一群组和第二群组中的至少一个。

[0010] 本发明可以提供改进的和/或便利的噪声衰减。在很多实施例中,对运算资源的需求大幅降低。在很多实施例中,该方法可以允许更有效的噪声衰减,其可以导致更快的噪声衰减。在很多场景中,该方法可以使得能够实现或者允许实时的噪声衰减。在很多场景和应用中,由于所考虑的可能候选的减少而导致的恰当码本条目的更精确估计,可以执行更精确的噪声衰减。

[0011] 每一个预期信号候选可以具有与该时间段的持续时间对应的持续时间。每一个噪声信号候选可以具有与该时间段的持续时间对应的持续时间。

[0012] 传感器信号可被分段到时间段中,这些时间段可以重叠,或者直接明确地对应于音频信号的时间段。在一些实施例中,分段器可以将传感器信号分段到与音频信号相同的时间段中。针对每个时间段的子集可以基于相同时间段中的传感器信号来确定。

[0013] 预期信号和噪声候选中的每一个都可以由表征信号分量的一组参数来表示。例如,每一个预期信号候选可以包括用于线性预测模型的一组线性预测系数。每一个预期信号候选可以包括表征频谱分布的一组参数,诸如例如功率谱密度(PSD)之类。

[0014] 噪声信号分量可以对应于并非预期信号分量的一部分的任何信号分量。例如,噪声信号分量可以包含白噪声、有色噪声、来自不想要的噪声源的确定性噪声等等。噪声信号分量可以是在不同时间段可能发生改变的非稳态噪声。对于每个时间段,噪声衰减器对每个时间段的处理可以是独立的。由此,音频环境中的噪声可以源自离散的声源,或者可以例如是回响或扩散声音分量。

[0015] 传感器信号可以接收自执行对预期来源和/或噪声的测量的传感器。

[0016] 子集可以分别属于第一和第二码本。特别地,当传感器信号提供对预期信号源的

测量时,该子集可以是第一码本的子集。当传感器信号提供对噪声的测量时,该子集可以是第二码本的子集。

[0017] 噪声估计器可被设置成为预期信号候选和噪声候选生成作为预期信号候选和噪声候选的加权组合、特别是加权总和的估计信号候选,其中权重被确定为使指示该时间段中的估计信号候选与音频信号之间的差的代价函数最小化。

[0018] 特别地,预期信号候选和/或噪声信号候选可以是可能信号分量的参数化表示。用于定义候选的参数数量可能典型地不超过20个,或者在很多实施例中有利地不超过10个。

[0019] 第一码本的预期信号候选和第二码本的噪声信号候选中的至少一个可以由频谱分布表示。特别地,候选可以由参数化的功率谱密度(PSD)的码本条目来表示,或者等价地由线性预测参数的码本条目来表示。

[0020] 在一些实施例中,与第一信号相比,传感器信号可以具有更小的频率带宽。在一些实施例中,噪声衰减装置可以接收多个传感器信号,并且子集的生成可以基于该多个传感器信号。

[0021] 特别地,噪声衰减器可以包含用于通过为第一码本的码本条目的第一群组中的预期信号候选与第二码本的码本条目的第二群组中的噪声信号候选的每个配对生成组合信号来生成多个估计信号候选的处理器、电路、功能单元或工具;用于从估计信号候选中生成用于在该时间段中的第一信号的信号候选的处理器、电路、功能单元或工具;用于响应于该信号候选来衰减在该时间段中第一信号的噪声的处理器、电路、功能单元或工具;以及用于通过响应于参考信号来选择码本条目的子集而生成第一群组和第二群组中的至少一个的处理器、电路、功能单元或工具。

[0022] 特别地,信号可以是音频信号,环境可以是音频环境,预期来源可以是音频源,并且噪声可以是音频噪声。

[0023] 特别地,噪声衰减装置可以包括:接收机,其用于接收针对音频环境的音频信号,所述音频信号包括对应于来自音频环境中的预期音频源的音频的预期信号分量,以及对应于音频环境中的噪声的噪声信号分量;第一码本,其包括用于预期信号分量的多个预期信号候选,每一个预期信号候选都表示可能的预期信号分量;第二码本,其包括用于噪声信号分量的多个噪声信号候选,每一个预期信号候选都表示可能的噪声信号分量;输入,其用于接收提供对音频环境的测量的传感器信号,所述传感器信号表示对音频环境中的预期音频源或噪声的测量;分段器,其用于将音频信号分段到时间段中;噪声衰减器,其被设置成为每一个时间段执行以下步骤:通过为第一码本的码本条目的第一群组中的预期信号候选与第二码本的码本条目的第二群组中的噪声信号候选的每个配对生成组合信号来生成多个估计信号候选;从估计信号候选中生成用于在该时间段中的音频信号的信号候选,以及响应于该信号候选来衰减在该时间段中的音频信号的噪声,其中所述噪声衰减器被设置成通过响应于参考信号来选择码本条目的子集而生成第一群组和第二群组中的至少一个。

[0024] 特别地,预期信号分量可以是语音信号分量。

[0025] 传感器信号可以接收自执行对预期来源和/或噪声的测量的传感器。该测量可以是例如通过一个或多个麦克风的声学测量,但其并不必如此。例如,在一些实施例中,该测量可以是机械或视觉测量。

[0026] 根据本发明的可选特征,传感器信号表示对预期来源的测量,并且噪声衰减器被

设置成通过从第一码本中选择码本条目的子集来生成第一群组。

[0027] 在很多实施例中,这样做可以允许降低的复杂度、便利的操作和/或改善的性能。在很多实施例中,可以为预期信号源生成特别有用的传感器信号从而允许要搜索的预期信号候选的数量的可靠降低。例如,对于预期信号源为语音源的情况,可以从骨传导麦克风生成语音信号的精确但却不同的表示。由此,在很多场景中可以有利的采用预期信号源的特定特性来提供基于与音频信号不同的传感器信号的潜在候选的大幅减少。

[0028] 根据本发明的可选特征,第一信号是音频信号,预期来源是音频源,预期信号分量是语音信号,并且传感器信号是骨传导麦克风信号。

[0029] 这样做可以提供特别有效且高性能的语音增强。

[0030] 根据本发明的可选特征,与预期信号分量相比,传感器信号提供精度较低的预期来源表示。

[0031] 本发明可以允许使用由质量降低的信号(并且因此潜在地不适于直接的信号衰减或信号再现)提供的附加信息来执行高质量的噪声衰减。

[0032] 根据本发明的可选特征,传感器信号表示对噪声的测量,并且噪声衰减器被设置成通过从第二码本选择码本条目的子集来生成第二群组。

[0033] 在很多实施例中,这样做可以允许降低的复杂度、便利的操作和/或改善的性能。在很多实施例中,可以为一个或多个噪声源(包括扩散噪声)生成特别有用的传感器信号从而允许要搜索的噪声信号候选的数量的可靠降低。在很多实施例中,与预期信号分量相比,噪声更为多变。例如,语音增强可被用在很多不同的环境中,并且因此可以在很多不同的噪声环境中使用。因此,在不同的环境中,噪声的特性可能会大幅变化,而语音特性则往往相对恒定。因此,噪声码本通常可以包含用于许多非常不同的环境的条目,并且在很多场景中,传感器信号可以允许生成对应于当前噪声环境的子集。

[0034] 根据本发明的可选特征,传感器信号是机械振动检测信号。

[0035] 在很多场景中,这样做可以允许特别可靠的性能。

[0036] 根据本发明的可选特征,传感器信号是加速度计信号。

[0037] 在很多场景中,这样做可以允许特别可靠的性能。

[0038] 根据本发明的可选特征,噪声衰减装置还包括映射器,其用于生成多个传感器信号候选与第一码本和第二码本中的至少一个的码本条目之间的映射;并且其中噪声衰减器被设置成响应于所述映射来选择码本条目的子集。

[0039] 在很多实施例中,这样做可以允许降低的复杂度、便利的操作和/或改善的性能。特别地,其可以允许便利和/或改进的适当候选子集的生成。

[0040] 根据本发明的可选特征,噪声衰减器被设置成响应于多个传感器信号候选中的每一个与传感器信号之间的距离量度来从多个传感器信号候选中选择第一传感器信号候选,并且响应于针对第一信号候选的映射来生成子集。

[0041] 在很多实施例中,这样做可以提供特别有利且切实可行的适当映射信息的生成,从而允许适当的候选子集的可靠生成。

[0042] 根据本发明的可选特征,映射器被设置成基于来自发出第一信号的输入传感器和发出传感器信号的传感器的同时测量来生成映射。

[0043] 这样做可以提供特别有效的实现,并且可以特别地降低复杂度,并且例如允许可

靠映射的便利和/或改善的确定。

[0044] 根据本发明的可选特征,映射器被设置成基于传感器信号候选与第一码本和第二码本中的至少一个的码本条目之间的差量度来生成映射。

[0045] 这样做可以提供特别有效的实现,并且可以特别地降低复杂度,并且例如允许可靠映射的便利和/或改善的确定。

[0046] 根据本发明的可选特征,第一信号是来自第一麦克风的麦克风信号,并且传感器信号是来自远离第一麦克风的第二麦克风的麦克风信号。

[0047] 在很多实施例中,这样做可以允许降低的复杂度、便利的操作和/或改善的性能。

[0048] 根据本发明的可选特征,第一信号是音频信号,并且传感器信号来自非音频传感器。

[0049] 在很多实施例中,这样做可以允许降低的复杂度、便利的操作和/或改善的性能。

[0050] 根据本发明的一个方面,提供了一种噪声衰减的方法,包括:接收针对环境的第一信号,所述第一信号包括对应于来自环境中的预期来源的信号的第一信号分量,以及对应于环境中的噪声的第二信号分量;提供包括用于预期信号分量的多个预期信号候选的第一码本,每一个预期信号候选都表示可能的预期信号分量;提供包括用于噪声信号分量的多个噪声信号候选的第二码本,每一个噪声信号候选都表示可能的噪声信号分量;接收提供对环境的测量的传感器信号,该传感器信号表示对环境中的预期来源或噪声的测量;将第一信号分段到时间段中;对于每一个时间段,执行以下步骤:通过为第一码本的码本条目的第一群组中的预期信号候选与第二码本的码本条目的第二群组中的噪声信号候选的每个配对生成组合信号来生成多个估计信号候选,从估计信号候选中生成用于在该时间段中的第一信号的信号候选,以及响应于该信号候选来衰减在该时间段中的第一信号的噪声;以及通过响应于参考信号来选择码本条目的子集而生成第一群组和第二群组中的至少一个。

[0051] 从以下所描述的一个或多个实施例并通过对其进行参考的阐述,本发明的这些和其他方面、特征和优点将是显而易见的。

附图说明

[0052] 参考附图,本发明的实施例将仅作为示例而被描述,在附图中

[0053] 图1是根据本发明的一些实施例的噪声衰减装置的元件示例的图示;

[0054] 图2是用于图1的噪声衰减装置的噪声衰减器的元件示例的图示;

[0055] 图3是根据本发明的一些实施例的噪声衰减装置的元件示例的图示;以及

[0056] 图4是根据本发明的一些实施例的用于噪声衰减装置的码本映射的图示。

具体实施方式

[0057] 以下描述集中在适用于音频噪声衰减、特别是适用于通过噪声衰减的语音增强的发明实施例上。然而将领会到,本发明并不局限于该应用,而是可以应用于很多其他信号。

[0058] 图1图示了根据本发明的一些实施例的噪声衰减器的示例。

[0059] 该噪声衰减器包括接收包含预期分量和非预期分量二者的信号的接收机101。所述非预期分量被称为噪声信号,并且可以包含并非预期信号分量的一部分的任何信号分量。预期信号分量对应于从预期声源生成的声音,而非预期或噪声信号分量则可以对应于

来自包含扩散和回响噪声等在内的所有其他声源的贡献。所述噪声信号分量可以包含环境中的周围环境噪声、来自非预期声源的音频等等。

[0060] 在图1的系统中,信号是音频信号,特别地,该信号可以从捕获给定音频环境中的音频信号的麦克风信号中生成。以下描述将集中在其中预期信号分量是来自预期说话者的语音信号的实施例上。

[0061] 接收机101耦合到将音频信号分段到时间段中的分段器103。在一些实施例中,所述时间段可以是不重叠的,但在其他实施例中,所述时间段可以是重叠的。另外,分段可以通过应用适当成形的窗口函数来执行,并且特别地,噪声衰减装置可以采用众所周知的使用诸如汉宁或汉明窗之类的适当窗口的分段的重叠和添加技术。时间段的持续时间将取决于具体实现,但是在很多实施例中,其将大约为10-100毫秒(msec)。

[0062] 分段器103被馈送至噪声衰减器105,所述噪声衰减器105执行基于分段的噪声衰减以突出与非预期噪声信号分量相对的预期信号分量。所得到的噪声衰减的分段被馈送至输出处理器107,所述输出处理器107提供连续的音频信号。特别地,输出处理器107可以例如通过执行重叠和添加功能来执行去分段。将领会的是,在其他实施例、例如在其中对噪声衰减的信号执行另外的基于分段的信号处理的实施例中,输出信号可以作为分段信号来提供。

[0063] 噪声衰减是基于码本方法,该方法使用涉及预期信号分量和涉及噪声信号分量的单独的码本。因此,噪声衰减器105耦合到第一码本109,所述第一码本109是预期信号码本,并且在该具体示例中是语音码本。噪声衰减器105还耦合到第二码本111,所述第二码本111是噪声信号码本。

[0064] 噪声衰减器105被设置成选择语音码本和噪声码本的码本条目,使得对应于所选择的条目的信号分量组合最近似于该时间段中的音频信号。一旦找到恰当的码本条目(连同其缩放),它们就表示所捕获的音频信号中的单独的语音信号分量和噪声信号分量的估计。特别地,对应于所选择的语音码本条目的信号分量是所捕获的音频信号中的语音信号分量的估计,并且噪声码本条目提供噪声信号分量的估计。因此,该方法使用了码本方法来估计音频信号的语音和噪声信号分量,并且一旦确定了这些估计,它们就可以被用来衰减音频信号中与语音信号分量相对的噪声信号分量,这是因为所述估计使得在这些之间做出区分成为可能。

[0065] 因此,在图1的系统中,噪声衰减器105耦合到预期信号码本109,其包括许多码本条目,其中每个码本条目都包括一组定义了可能的预期信号分量的参数,并且在该具体示例中,所述预期信号分量是预期语音信号。类似地,噪声衰减器105耦合到噪声信号码本109,其包括许多码本条目,其中每个码本条目都包括一组定义了可能的噪声信号分量的参数。

[0066] 针对预期信号分量的码本条目对应于用于预期信号分量的潜在候选,并且针对噪声信号分量的码本条目对应于用于噪声信号分量的潜在候选。每一个条目都包括一组分别表征可能的预期信号或噪声分量的参数。在该具体示例中,第一码本109的每个条目都包括一组表征可能的语音信号分量的参数。因此,由该码本中的码本条目表征的信号具有语音信号特性,并且因而码本条目将语音特性的认识引入到语音信号分量的估计中。

[0067] 针对预期信号分量的码本条目可以基于预期音频源的模型,或者可以附加地或可

替换地由训练过程来确定。例如,码本条目可以是用于为了表示语音特性而被开发的语音模型的参数。作为另一示例,可以记录和统计处理大量的语音样本以生成存储在码本中的适当数量的潜在语音候选。类似地,针对噪声信号分量的码本条目可以基于噪声模型,或者可以附加地或可替换地由训练过程来确定。

[0068] 特别地,码本条目可以基于线性预测模型。实际上,在该具体示例中,码本的每个条目都包括一组线性预测参数。特别地,码本条目通过训练过程生成,其中线性预测参数通过与大量的信号样本相适应而生成。

[0069] 在一些实施例中,码本条目可被表示为频率分布,并且特别地被表示为功率谱密度(PSD)。PSD可以直接对应于线性预测参数。

[0070] 典型地,用于每个码本条目的参数的数量相对较小。实际上,规定每个码本条目的参数典型地不超过20个,并且通常不会超过10个。由此,使用了预期信号分量的相对粗略的估计。这样做允许降低的复杂度和便利的处理,但是已证实其在大多数情况下仍提供有效的噪声衰减。

[0071] 更详细来说,考虑其中语音与噪声被假设为独立的相加噪声模型:

$$[0072] \quad y(n) = x(n) + w(n),$$

[0073] 其中 $y(n)$ 、 $x(n)$ 和 $w(n)$ 分别表示被采样的含噪语音(输入音频信号)、纯净语音(预期语音信号分量)以及噪声(噪声信号分量)。

[0074] 基于码本的噪声衰减典型地包含将码本搜索遍以找到分别针对信号分量和噪声分量的码本条目,使得缩放组合与所捕获的信号最为相似,从而为每个短时间段提供语音和噪声分量的估计。设 $P_y(\omega)$ 表示观测的含噪信号 $y(n)$ 的功率谱密度(PSD), $P_x(\omega)$ 表示语音信号分量 $x(n)$ 的PSD,并且 $P_w(\omega)$ 表示噪声信号分量 $w(n)$ 的PSD,则

$$[0075] \quad P_y(\omega) = P_x(\omega) + P_w(\omega)$$

[0076] 设 $\hat{\cdot}$ 表示对应的PSD的估计,基于传统码本的噪声衰减可以通过将频域维纳(Wiener)滤波器 $H(\omega)$ 应用到所捕获的信号来减小噪声,即:

$$[0077] \quad P_m(\omega) = \hat{P}_y(\omega) H(\omega)$$

[0078] 其中维纳滤波器由下式给出:

$$[0079] \quad H(\omega) = \frac{\hat{P}_x(\omega)}{\hat{P}_x(\omega) + \hat{P}_w(\omega)},$$

[0080] 码本分别包括语音信号候选和噪声信号候选,并且关键的问题是标识出最适合的候选配对以及每一个的相对加权。

[0081] 语音和噪声PSD的估计和由此的恰当候选的选择可以遵循最大似然(ML)方法或贝叶斯最小均方误差(MMSE)方法。

[0082] 线性预测系数的矢量与基础PSD之间的关系可以由下式来确定

$$[0083] \quad P_x(\omega) = \frac{1}{|A_x(\omega)|^2},$$

[0084] 其中 $\theta_x = (a_{x_0}, \dots, a_{x_p})$ 是线性预测系数, $a_{x_0} = 1$ 和 p 是线性预测模型阶数,并且

$$A_x(\omega) = \sum_{k=0}^p a_{x,k} e^{-j\omega k}.$$

[0085] 通过使用该关系,可以由下式给出所捕获的信号估计PSD

$$[0086] \quad \hat{P}_y(\omega) = \underbrace{g_x P_x(\omega)}_{\approx \hat{P}_x(\omega)} + \underbrace{g_w P_w(\omega)}_{\approx \hat{P}_w(\omega)},$$

[0087] 其中 g_x 和 g_w 是与语音和噪声PSD相关联的频率无关的水平增益。引入这些增益以顾及存储在码本中的PSD与在输入音频信号中遇到的PSD之间的水平变化。

[0088] 常规方法是基于搜索遍语音码本条目与噪声码本条目的所有可能配对来确定使观测的含噪PSD与估计PSD之间的某个相似性量度最大化的配对,正如下文所描述的那样。

[0089] 考虑由来自语音码本的第 i 个PSD和来自噪声码本的第 j 个PSD给出的语音和噪声PSD的配对。对应于该配对的含噪PSD可被写为

$$[0090] \quad \hat{P}_y^j(\omega) = g_x^j P_x^i(\omega) + g_w^j P_w^j(\omega).$$

[0091] 在该方程中,PSD是已知的,而增益是未知的。因此对于语音和噪声PSD的每一个可能的配对而言,必须确定增益。这可以基于最大似然方法来完成。预期语音和噪声PSD的最大似然估计可以在两步过程中获取。已经导致观测的含噪PSD的给定配对 $g_x^j P_x^i(\omega)$ 和 $g_w^j P_w^j(\omega)$ 的似然的对数由以下方程表示:

$$[0092] \quad L_y(P_y(\omega), \hat{P}_y^j(\omega)) = \int_0^{2\pi} -\frac{P_y(\omega)}{\hat{P}_y^j(\omega)} + \ln\left(\frac{1}{\hat{P}_y^j(\omega)}\right) d\omega$$

$$= \int_0^{2\pi} -\frac{P_y(\omega)}{g_x^j P_x^i(\omega) + g_w^j P_w^j(\omega)}$$

$$[0093] \quad + \ln\left(\frac{1}{g_x^j P_x^i(\omega) + g_w^j P_w^j(\omega)}\right) d\omega.$$

[0094] 在第一步骤中,确定使 $L_y(P_y(\omega), \hat{P}_y^j(\omega))$ 最大化的未知水平项 g_x^j 和 g_w^j 。做到这一点的一种方式是通过区分关于 g_x^j 和 g_w^j , 将结果设置成0,并且解出联立方程的结果集。然而,这些方程是非线性且不服从闭合形式解的。可替换的方法是基于这样的事实,即似然在 $P_y(\omega) = \hat{P}_y^j(\omega)$ 时被最大化,并且因此可以通过最小化这两个实体之间的频谱距离来获取增益项。

[0095] 一旦已知水平项,则由于所有实体都是已知的,就能够确定 $L_y(P_y(\omega), \hat{P}_y^j(\omega))$ 的值。针对语音和噪声码本条目的所有配对重复地执行该过程,并且使用导致最大的似然的配对来获取语音和噪声PSD。由于为每一个短时间段执行了该步骤,因此即便在非稳态噪

声条件下,该方法也能够精确估计噪声PSD。

[0096] 设 $\{i^*, j^*\}$ 表示导致针对给定分段的最大的似然的配对,并且设 g_x^* 和 g_w^* 表示对应的水平项。则语音和噪声PSD由下式给出

$$[0097] \quad \hat{P}_x(\omega) = g_x^* P_x^j$$

$$[0098] \quad \hat{P}_w(\omega) = g_w^* P_w^j$$

[0099] 由此,这些结果定义了被应用于输入音频信号以生成噪声衰减的信号的维纳滤波器。

[0100] 因此,现有技术是基于找到作为针对语音信号分量的良好估计的适当预期信号码本条目和作为针对噪声信号分量的良好估计的适当噪声信号码本条目。一旦找到这些条目,就可以应用有效的噪声衰减。

[0101] 然而,该方法非常复杂并且对资源要求高。特别地,为了找到最佳匹配,必须评估噪声与语音码本条目的所有可能配对。另外,由于码本条目必须表示各种各样的可能信号,因此这会导致非常大的码本,并且由此导致必须评估的很多可能配对。特别地,噪声信号分量在可能的特性中可能通常具有例如取决于具体的使用环境等等的大的变化。因此,为了确保足够接近的估计,通常需要非常大的噪声码本。这会导致非常高的运算需求。

[0102] 在图1的系统中,通过使用第二信号来减少噪声衰减算法搜索的码本条目的数量,可以大幅降低该算法的复杂度,特别是该算法的运算资源使用。特别地,除了从麦克风接收用于噪声衰减的音频信号之外,该系统还接收传感器信号,所述传感器信号提供主要是预期信号分量或者主要是噪声信号分量的测量。

[0103] 因此,图1的噪声衰减器包括从适当传感器接收传感器信号的传感器接收器113。该传感器信号提供对音频环境的测量,使得其表示对预期音频源的测量或对音频环境的测量。

[0104] 在该示例中,传感器接收器113耦合到分段器103,该分段器103着手将传感器信号分段到与音频信号相同的时间段中。然而将领会到,这种分段是可选的,并且在其他实施例中,传感器信号可以例如被分段到相对于音频信号的分段更长、更短、重叠、不相交等的时间段中。

[0105] 因此,在图1的示例中,噪声衰减器105为每个分段接收音频信号和传感器信号,其中所述传感器信号提供对音频环境中的预期音频源或噪声的不同测量。然后,噪声衰减器使用由传感器信号提供的附加信息来为对应码本选择码本条目的子集。由此,当传感器信号表示对预期音频源的测量时,噪声衰减器105生成预期信号候选的子集。然后,在噪声码本111中的噪声信号候选与所生成的预期信号候选子集中的候选的可能配对之上执行搜索。当传感器信号表示对噪声环境的测量时,噪声衰减器105从噪声码本111生成预期噪声候选的子集。然后,在预期信号码本109中的预期信号候选与所生成的噪声信号候选子集中的候选的可能配对之上执行搜索。

[0106] 图2图示了噪声衰减器105的一些元件的示例。该噪声衰减器包括估计处理器201,其通过为预期信号码本的码本条目的第一群组中的预期信号候选与噪声码本的码本条目的第二群组中的噪声信号候选的每个配对生成组合信号来生成多个估计信号候选。由此,

估计处理器201为来自噪声码本的候选(码本条目)的群组的噪声候选与来自预期信号码本的候选(码本条目)的群组的预期信号候选的每个配对生成接收信号的估计。特别地,针对候选配对的估计可以作为加权和、并且特别地作为导致代价函数最小化的加权总和而生成。

[0107] 噪声衰减器105还包括群组处理器203,其被设置成通过响应于参考信号来选择码本条目的子集而生成第一群组和第二群组中的至少一个。由此,第一群组或第二群组可以简单地等同于整个码本,但是群组中的至少一个是作为码本的子集生成的,其中所述子集在传感器信号的基础上生成。

[0108] 估计处理器201还耦合到候选处理器205,所述候选处理器205着手从估计信号候选中生成用于该时间段中的输入信号的信号候选。例如,该候选可以简单地通过选择导致最低代价函数的估计来生成。可替换地,该候选可以作为估计的加权组合来生成,其中权重取决于代价函数的值。

[0109] 候选处理器205耦合到噪声衰减处理器207,所述噪声衰减处理器207着手响应于所生成的信号候选来衰减该时间段中的输入信号的噪声。例如,可以如先前所描述的那样应用维纳滤波器。

[0110] 由此,可以使用第二传感器信号来提供可以被用于控制搜索的附加信息,使得搜索可以大幅减少。然而,传感器信号不直接影响音频信号,而是仅仅引导搜索来找到最优估计。结果,由传感器进行的测量中的失真、噪声、不精确性等等将不直接影响信号处理或噪声衰减,并且因此将不直接引入任何信号质量降级。这样一来,传感器信号可以具有大幅降低的质量,并且特别是对预期信号测量而言,如果直接被使用,传感器信号可能是将会提供不适宜的音频(尤其是语音)质量的信号。这样一来,可以使用多种多样的传感器,特别是可以提供与捕获音频信号的麦克风截然不同的信息的传感器,诸如例如非音频传感器之类。

[0111] 在一些实施例中,传感器信号可以表示对预期音频源的测量,其中与音频信号的预期信号分量相比,传感器信号特别地提供精度较低的预期音频源表示。

[0112] 例如,可以使用麦克风来从嘈杂环境中的人捕获语音。可以使用不同类型的传感器以提供对语音信号的不同测量,尽管该语音信号的质量可能不足以提供可靠的语音但仍可用于收窄在语音码本中的搜索。

[0113] 主要地仅捕获预期信号的参考传感器的示例为可以被佩戴在用户咽喉附近的骨传导麦克风。这种骨传导麦克风将捕获通过(人体)组织传播的语音信号。由于该传感器与用户身体接触并且从外部声学环境屏蔽开,因此其能够捕获具有非常高的信噪比的语音信号,即,其提供以骨传导麦克风信号形式的传感器信号,其中由预期音频源(说话者)引起的信号能量大大高于(比方说高于至少10dB或更多)由其他来源引起的信号能量。

[0114] 然而,由于传感器的位置,所捕获的信号的质量远不同于由放置在用户嘴部前方的麦克风拾取的空气传导语音的质量。因此,所得到的质量不足以被直接用作语音信号,但是非常适于引导基于码本的噪声衰减以仅仅搜索语音码本的小型子集。

[0115] 因此,与需要使用大型语音和噪声码本的联合增强的常规方法不同,由于纯净的参考信号的存在,图1的方法只需要在语音码本的小型子集之上执行最优化。由于可能的组合的数量随候选数量的降低而急剧降低,因此这导致在运算复杂度中的显著节约。此外,纯净的参考信号的使用使得能够实现对真实纯净语音(即预期信号分量)进行精密建模的语

音码本子集的选择。因此,选择错误候选的可能性大幅降低,并且由此可以改善整个噪声衰减的性能。

[0116] 在其他实施例中,传感器信号可以表示对音频环境中的噪声的测量,并且噪声衰减器105可被设置成降低所考虑的噪声码本111中的候选/条目的数量。

[0117] 噪声测量可以是对音频环境的直接测量,或者可以例如是使用不同模态的传感器(即使用非音频传感器)的间接测量。

[0118] 作为音频传感器的示例,其可以是远离捕获音频信号的麦克风而定位的麦克风。例如,捕获语音信号的麦克风可以靠近说话者的嘴部而定位,而第二麦克风用于提供传感器信号。第二麦克风可以定位在其中噪声凌驾于语音信号之上的定位处,并且特别地可以定位得足够远离说话者的嘴部。音频传感器可以足够远,以使源自预期声源的能量与噪声能量之间的比与在所捕获的音频信号中相比,在传感器信号中降低不少于10dB。

[0119] 在一些实施例中可以使用非音频传感器来生成例如机械振动检测信号。例如,可以使用加速度计来生成以加速度计信号形式的传感器信号。这样的传感器可以例如被安装在通信设备上并检测其振动。作为另一示例,在其中已知特定机械实体是主要噪声源的实施例中,可以将加速度计附着到该设备,以提供非音频传感器信号。作为具体示例,在洗衣店应用中,加速度计可被定位在洗衣机或甩干机上。

[0120] 作为另一示例,传感器信号可以是视觉检测信号。例如,可以使用摄像机来检测指示音频环境的视觉环境特性。例如,视频检测可以允许对给定噪声源是否活动的检测,并且可以被用于将噪声候选的搜索缩减至对应的子集。(视觉传感器信号还可以用于降低所搜索的预期信号候选的数量,例如通过将唇读算法应用于说话人来得到适当候选的粗略指示,或者例如通过使用面部识别系统来检测说话者,使得可以选择对应的码本条目)。

[0121] 这样的噪声参考传感器信号然后可以被用于选择被搜索的噪声码本条目的子集。这样做不但可以有效降低必须考虑的码本条目配对的数量,并且因此大幅降低复杂度,而且还可以导致更精确的噪声估计以及因此的改进的噪声衰减。

[0122] 传感器信号表示对预期信号源或噪声的测量。然而将领会到,传感器信号还可以包含其他信号分量,并且特别地,在一些场景中,传感器信号可以包含来自预期声源和环境中的噪声二者的贡献。然而,在传感器信号中,这些分量的分布或权重将是不同的,并且特别地,分量中的一个将典型是主导的。典型地,与针对其确定子集的码本对应的分量(即预期信号或噪声信号)的能量/功率比其他分量的能量高不小于3dB、10dB或甚至20dB。

[0123] 一旦在码本条目的所有候选配对之上都执行了搜索,就为每个配对生成信号候选估计,典型地还会连带生成估计与被测音频信号适应得多么接近的指示。然后,基于估计信号候选来为该时间段生成信号候选。所述信号候选可以通过考虑导致所捕获的音频信号的信号候选的似然估计而被生成。

[0124] 作为低复杂度的示例,系统可以简单地选择具有最高似然值的估计信号候选。在更复杂的实施例中,信号候选可以通过所有估计信号候选的加权组合特别是总和来计算,其中每一个估计信号候选的权重取决于似然值的对数。

[0125] 然后,基于所计算的信号候选来补偿音频信号。特别地,通过用以下维纳滤波器对音频信号进行滤波:

$$[0126] \quad H(\omega) = \frac{\hat{P}_s(\omega)}{\hat{P}_s(\omega) + \hat{P}_n(\omega)}.$$

[0127] 将领会到,可以使用基于所估计的信号及噪声分量来减小噪声的其他方法。例如,该系统可以从输入音频信号中减去估计噪声候选。

[0128] 因此,噪声衰减器105从在其中噪声信号分量相对于语音信号分量被衰减的时间段中的输入信号生成输出信号。

[0129] 将领会到,在不同的实施例中可以使用不同方法来确定码本条目的子集。例如,在一些实施例中,传感器信号可被等价地参数化成码本条目,例如通过将其表示为具有对应于那些码本条目的参数的PSD(特别地,对每一个参数使用相同的频率范围)。然后,可以使用诸如平方误差之类的适当的距离量度来找到传感器信号PSD与码本条目之间的最接近的匹配。之后,噪声衰减器105可以选择最接近于所标识的匹配的预定数量的码本条目。

[0130] 然而,在很多实施例中,噪声衰减系统可以被设置成基于传感器信号候选与码本条目之间的映射来选择子集。由此,该系统可以包括如图2所示的映射器301,其中该映射器301被设置成生成从传感器信号候选到码本候选的映射。

[0131] 该映射从映射器301馈送至噪声衰减器105,并且在所述噪声衰减器105中被用于生成其中一个码本的子集。图3图示了噪声衰减器105如何可以操作于其中传感器信号针对预期信号的示例的一个示例。

[0132] 在该示例中,为所接收的传感器信号生成线性LPC参数,并且所得到的参数被量化以对应于所生成的映射401中的可能传感器信号候选。映射401提供从包括传感器信号候选的传感器信号码本到语音码本109中的语音信号候选的映射。该映射被用于生成语音码本条目的子集403。

[0133] 特别地,噪声衰减器105可以搜索遍映射401中的所存储的传感器信号候选,以根据适当的距离量度来确定最接近于被测传感器的传感器信号候选,所述距离量度例如是参数的误差平方和。然后,其可以基于子集例如通过将映射至所标识的传感器信号候选的一个或多个语音信号候选包含在该子集中来生成映射。可以例如通过包含到所选择的语音信号候选的给定距离量度小于给定阈值的所有语音信号候选,或者通过包含被映射至所选择的传感器信号候选的给定距离量度小于给定阈值的传感器信号候选的所有语音信号候选来将子集生成成为具有预期的大小。

[0134] 基于音频信号,在子集403以及噪声码本111的条目之上执行搜索,以生成估计信号候选,并且然后生成针对该分段的信号候选,如前文所描述的那样。将领会到,相同的方法可以可替换地或附加地基于噪声传感器信号而被应用于噪声码本111。

[0135] 特别地,映射可以由可以生成码本条目和传感器信号候选二者的训练过程来生成。

[0136] 用于特定信号的N条目码本的生成可以基于训练数据,并且可以例如基于Y. Linde、A. Buzo以及R. Gray在1980年1月发表于Communications, IEEE Transactions on (IEEE会报:通信)第28卷第1号第84-95页的“An algorithm for vector quantizer design(用于矢量量化器设计的算法)”中所描述的Linde-Buzo-Gray (LBG)算法。

[0137] 特别地, 设 X 表示 L 个长度为 M 的训练矢量的集合, 其中元素 $\mathbf{x}_k \in X (1 \leq k \leq L)$ 。该算法始于运算对应于训练矢量的均值的单个码本条目, 即 $\mathbf{c}_0 = \overline{X}$ 。然后, 该条目被一分为二, 使得

$$\begin{aligned} \mathbf{c}_1 &= (1+\eta)\mathbf{c}_0 \\ \mathbf{c}_2 &= (1-\eta)\mathbf{c}_0, \end{aligned} \quad [0138]$$

[0139] 其中 η 是小常量。然后, 该算法将训练矢量划分成两个分区 X_1 和 X_2 , 使得

$$\mathbf{x}_k \in \begin{cases} X_1 & \text{iff } d(\mathbf{x}_k, \mathbf{c}_1) < d(\mathbf{x}_k, \mathbf{c}_2) \\ X_2 & \text{iff } d(\mathbf{x}_k, \mathbf{c}_2) < d(\mathbf{x}_k, \mathbf{c}_1) \end{cases} \quad [0140]$$

[0141] 其中 $d(\cdot; \cdot)$ 是某个失真量度, 例如均方误差(MSE)或加权MSE(WMSE)。然后, 根据下式来重定义当前的码本条目:

$$\begin{aligned} \mathbf{c}_1 &= \overline{X_1} \\ \mathbf{c}_2 &= \overline{X_2}. \end{aligned} \quad [0142]$$

[0143] 重复前两个步骤直至总码本误差不随当前码本条目改变。然后, 每个码本条目再次被拆分, 并且重复相同的过程直至条目数量等于 N 。

[0144] 设 R 和 Z 分别表示针对由参考传感器和音频信号麦克风捕获的相同声源(预期或非预期/噪声)的训练矢量集合。基于这些训练矢量, 可以生成传感器信号候选与长度为 N_d 的主码本(术语“主”酌情表示噪声或预期码本)之间的映射。

[0145] 码本可以例如通过首先使用上述LBG算法独立地生成映射(即传感器候选与主候选)的两个码本, 并且随后创建这些码本的条目之间的映射而被生成。映射可以基于码本条目的所有配对之间的距离量度以便在传感器码本与主码本之间创建一对一(或一对多/多对一)的映射。

[0146] 作为另一示例, 针对传感器信号的码本生成可以与主码本一起生成。特别地, 在该示例中, 映射可以基于来自发出音频信号的麦克风和来自发出传感器信号的传感器的同时测量。由此映射是基于在相同时间处捕获相同音频环境的不同信号的。

[0147] 在这样的示例中, 映射可以基于信号在时间上同步的假设, 并且传感器候选码本可以通过使用由将LBG算法应用于主训练矢量所引起的最终分区而被得到。如果(主码本)分区的集合被给出为

$$\mathbf{Z} = \{\mathbf{Z}_1, \mathbf{Z}_2, \dots, \mathbf{Z}_{N_d}\}, \quad [0148]$$

[0149] 那么可以生成对应于参考传感器 R 的分区集合, 使得:

$$\mathbf{r}_k \in R_j \text{ iff } \mathbf{z}_k \in \mathbf{Z}_j \quad 1 \leq k \leq L, 1 \leq j \leq N_d. \quad [0150]$$

[0151] 然后可以如先前所描述的那样应用所得到的映射。

[0152] 该系统可被用在众多不同的应用中, 包括例如需要单个麦克风降噪的应用, 例如移动电话学和DECT电话。作为另一示例, 该方法可被用在多麦克风语音增强系统(例如助听器、基于阵列的免提系统等等)中, 这些系统常常具有用于进一步降噪的单通道后处理器。

[0153] 实际上,虽然先前的描述针对的是音频信号中的音频噪声的衰减,但是将领会的是,所描述的原理和方法可以适用于其他类型的信号。实际上,要指出的是,通过使用所描述的码本方法,可以对包括预期信号分量和噪声的任何输入信号进行噪声衰减。

[0154] 这样的非音频实施例的示例可以是其中使用加速度计做出呼吸率的测量的系统。在这种情形中,测量传感器可以被放置在被测人员的胸部附近。此外,一个或多个附加的加速度计可以被定位在足部(或双脚)上,以移除在行走/跑步期间可能出现在一个或多个主加速度计信号上的噪声贡献。因此,安装在测试人员足部上的这些加速度计可被用于收窄噪声码本搜索。

[0155] 还将领会的是,多个传感器和传感器信号可以被用于生成被搜索的码本条目的子集。这些多个传感器信号可以被单独使用,也可以并行使用。例如,所使用的传感器信号可以取决于信号的种类、类别或特性,并且因此可以使用准则来选择子集生成所基于的传感器信号。在其他示例中,可以使用更加复杂的准则或算法来生成子集,其中准则或算法同时考虑多个传感器信号。

[0156] 将领会到,为了清楚起见,以上描述参考了不同的功能电路、单元和处理器来描述本发明的实施例。然而将显而易见的是,在不脱离于本发明的情况下,不同的功能电路、单元或处理器之间的功能性的任何适当分布都是可以使用的。例如,被图示说明成由单独的处理器或控制器执行的功能可以由相同的处理器或控制器执行。因此,对特定功能单元或电路的引用应仅被视为对用于提供所描述的功能性的适当工具的引用,而不是指示严格的逻辑或物理结构或组织。

[0157] 本发明可以以包含硬件、软件、固件或其任何组合在内的任何适当形式来实现。本发明可以可选地至少部分地作为在一个或多个数据处理器和/或数字信号处理器上运行的计算机软件来实现。本发明的实施例的元件和组件可以以任何适当的方式被物理地、功能性地和逻辑地实现。实际上,功能性可以实现在单个单元中、多个单元中,或者作为其他功能单元的一部分来实现。就此而论,本发明可以实现在单个单元中,或者可以物理地和功能性地分布在不同的单元、电路和处理器之间。

[0158] 虽然已经结合一些实施例来描述了本发明,但这并不旨在将其局限于这里所阐述的特定形式。相反,本发明的范围仅仅由随附的权利要求限制。此外,虽然特征可以表现为结合特定实施例而被描述,但是本领域技术人员将会认识到,所描述的实施例的各种特征可以根据本发明进行组合。在权利要求中,术语包括不排除其他元件或步骤的存在。

[0159] 另外,虽然被单独列出,但是多个工具、元件、电路或方法步骤可以由例如单个电路、单元或处理器来实现。此外,虽然可以将各个特征包含在不同的权利要求中,但是这些特征可以可能地被有利地组合,并且在不同权利要求中的包含并不暗示特征的组合不可行和/或不是有利的。另外,特征在一种权利要求类别中的包含并不暗示对该类别的限制,相反,这指示该特征同样可以酌情适用于其他的权利要求类别。此外,特征在权利要求中的顺序并不暗示特征必须以任何特定的顺序工作,并且特别地,在方法权利要求中的各个步骤的顺序并不暗示必须以该顺序执行这些步骤。相反,这些步骤可以以任何适当的顺序执行。另外,单数的引用并不排除多个。因此,对“一”、“一个”、“第一”、“第二”等等的引用并不排除多个。权利要求中的参考标记仅仅作为澄清示例而被提供,而不应将其解释为以任何方式限制权利要求的范围。

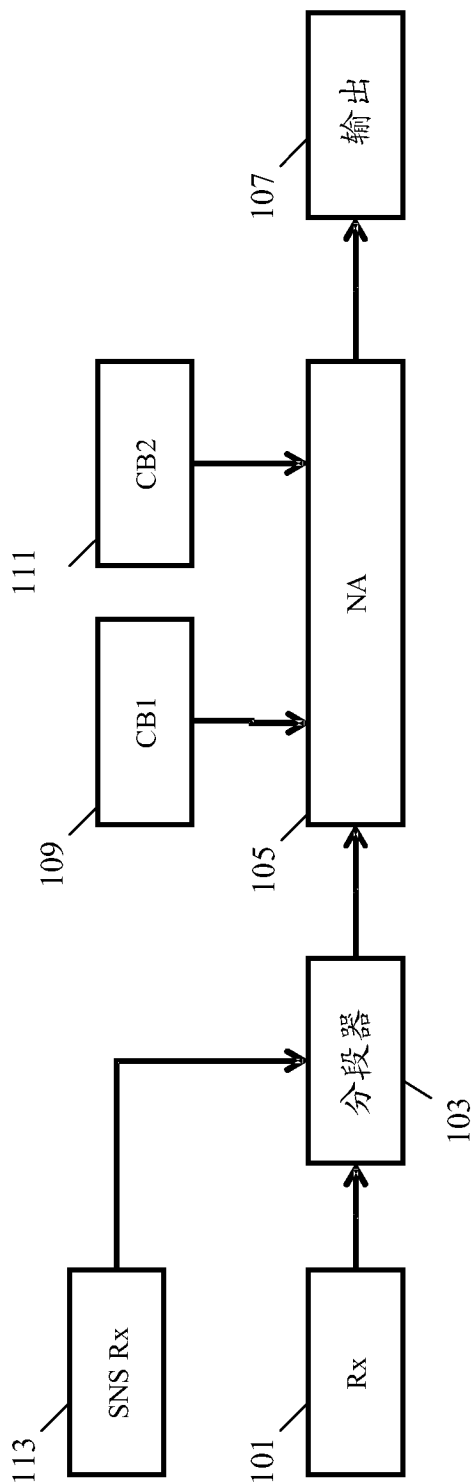


图 1

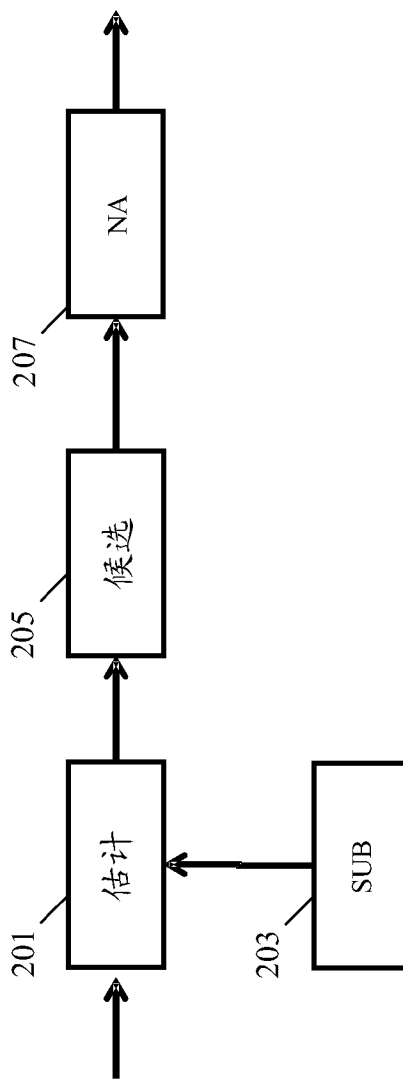


图 2

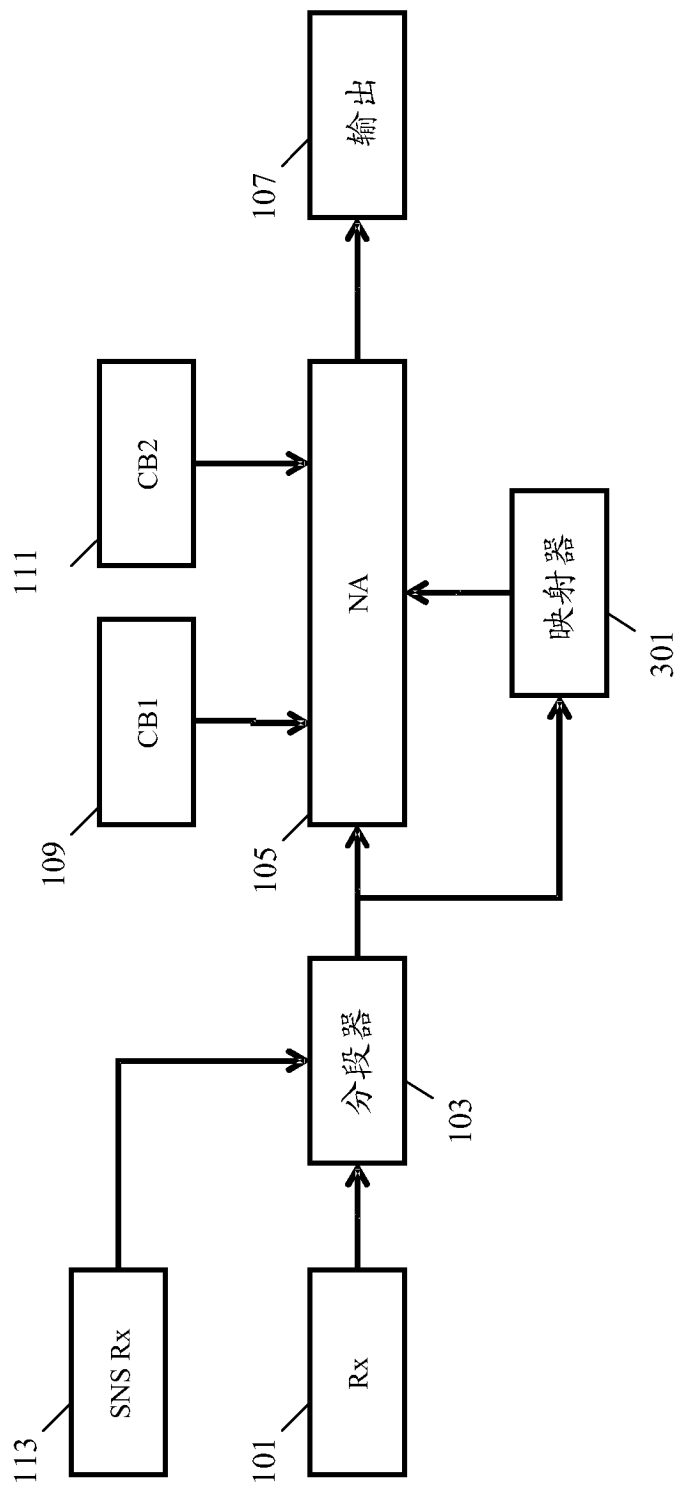


图 3

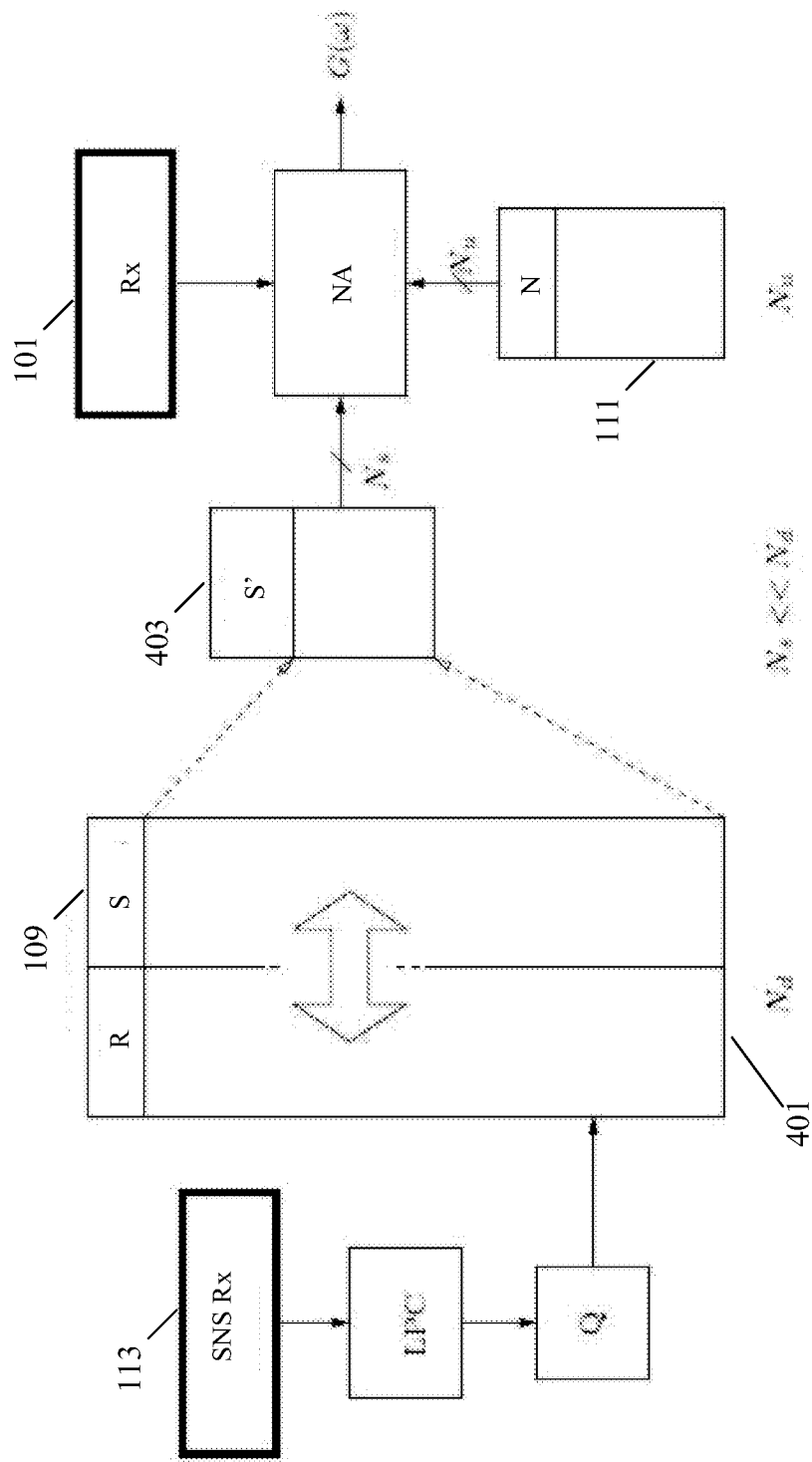


图 4