



## (12) 发明专利

(10) 授权公告号 CN 102111706 B

(45) 授权公告日 2015. 02. 25

(21) 申请号 201010621662. 2

(22) 申请日 2010. 12. 29

(30) 优先权数据

09180883. 2 2009. 12. 29 EP

(73) 专利权人 GN 瑞声达 A/S

地址 丹麦巴勒鲁普

(72) 发明人 卡尔 - 弗雷德里克 · 约翰 · 格兰

(74) 专利代理机构 中原信达知识产权代理有限  
责任公司 11219

代理人 谷惠敏 穆德骏

(51) Int. Cl.

H04R 25/00(2006. 01)

(56) 对比文件

US 2008253596 A1, 2008. 10. 16,

US 2008107297 A1, 2008. 05. 08,

US 7206421 B1, 2007. 04. 17,

W0 2007128825 A1, 2007. 11. 15,

CN 101529929 A, 2009. 09. 09,

US 2008253596 A1, 2008. 10. 16,

CN 101356854 A, 2009. 01. 28,

审查员 赵伟

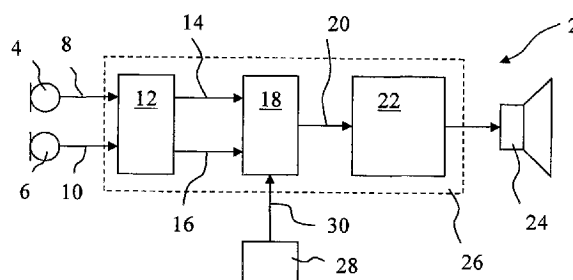
权利要求书2页 说明书14页 附图10页

(54) 发明名称

助听器中的波束形成

(57) 摘要

助听器中的波束形成。本发明涉及一种具有执行自适应双耳波束形成的能力的助听器系统,包括:第一麦克风和第二麦克风,用于提供电子输入信号;波束形成器,用于至少部分地基于所述电子输入信号来提供具有定向空间特性的第一音频信号,其中,所述波束形成器进一步被配置成至少部分地基于所述电子输入信号来提供第二音频信号,所述第二音频信号具有与所述第一音频信号不同的另一个空间特性,所述助听器系统进一步包括混合器,所述混合器被配置用于混合所述第一音频信号和第二音频信号,以便提供要由用户听到的输出信号。



1. 一种助听器系统,包括:

第一麦克风和第二麦克风,所述第一麦克风和第二麦克风用于提供电子输入信号;

波束形成器,所述波束形成器用于至少部分地基于所述电子输入信号来提供具有定向空间特性的第一音频信号,其特征在于,所述波束形成器进一步被配置成至少部分地基于所述电子输入信号来提供第二音频信号,所述第二音频信号具有与所述第一音频信号不同的另一个空间特性,

所述助听器系统进一步包括

混合器,所述混合器被配置用于混合所述第一音频信号和第二音频信号,以便提供要由用户听到的输出信号,

其中,所述波束形成器是自适应的。

2. 根据权利要求1所述的助听器系统,进一步包括处理器,该处理器被配置成根据听力损害校正算法来处理所述混合信号。

3. 根据权利要求1所述的助听器系统,进一步包括处理器,该处理器被配置成在混合所述第一音频信号和第二音频信号之前,根据听力损害校正算法来处理所述第一音频信号。

4. 根据权利要求1或2所述的助听器系统,其中,所述助听器系统包括用户操作的界面,所述用户操作的界面可操作地连接到所述混合器以用于控制所述第一音频信号和第二音频信号的混合。

5. 根据权利要求4所述的助听器系统,其中,所述用户操作的界面被置于单独的遥控设备中,所述单独的遥控设备经由无线链路可操作地连接到所述混合器。

6. 根据权利要求4所述的助听器系统,其中,所述用户操作的界面包括手动可操作开关。

7. 根据权利要求1或2所述的助听器系统,其中,所述助听器系统是双耳助听器系统,所述双耳助听器系统包括经由通信链路而彼此互连的第一助听器和第二助听器,并且其中,所述第一麦克风位于所述第一助听器中,并且所述第二麦克风位于所述第二助听器中。

8. 根据权利要求7所述的助听器系统,其中,所述第一助听器和第二助听器的每一个包括连接到所述波束形成器的附加麦克风。

9. 根据权利要求4所述的助听器系统,其中,所述助听器系统是双耳助听器系统,所述双耳助听器系统包括经由通信链路而彼此互连的第一助听器和第二助听器,并且其中,所述第一麦克风位于所述第一助听器中,并且所述第二麦克风位于所述第二助听器中。

10. 根据权利要求6所述的助听器系统,其中,所述助听器系统是双耳助听器系统,所述双耳助听器系统包括经由通信链路而彼此互连的第一助听器和第二助听器,所述第一麦克风位于所述第一助听器中,并且所述第二麦克风位于所述第二助听器中,并且其中,所述手动可操作开关被置于所述第一助听器和/或第二助听器中。

11. 根据权利要求1或2所述的助听器系统,形成双耳助听器系统的一部分。

12. 根据权利要求1或2所述的助听器系统,其中,所述第一音频信号和第二音频信号的空间特性是基本上互补的。

13. 根据权利要求1或2所述的助听器系统,其中,所述第二音频信号的所述空间特性是基本上全向的。

14. 根据权利要求 1 或 2 所述的助听器系统,其中,所述第一音频信号和第二音频信号的空间特性以如下这样的方式被产生:结果产生的所述混合音频信号的空间特性是基本上全向的。

15. 根据权利要求 1 或 2 所述的助听器系统,其中,所述助听器系统包括分类器,所述分类器用于提供周围的声音环境的分类,并且其中,所述助听器系统被配置用于根据所述分类来执行所述第一音频信号和第二音频信号的混合。

16. 根据权利要求 15 所述的助听器系统,其中,所述助听器系统被配置用于自动执行所述混合,使得能够对于不同的收听情况优化所述混合。

17. 根据权利要求 4 所述的助听器系统,其中,所述助听器系统包括分类器,所述分类器用于提供周围的声音环境的分类,所述助听器系统被配置用于根据所述分类来执行所述第一音频信号和第二音频信号的混合,所述助听器系统被配置用于自动执行所述混合,使得能够对于不同的收听情况优化所述混合,并且其中所述用户能够通过激活所述用户操作的界面来否决所述自动混合。

18. 根据权利要求 7 所述的助听器系统,其中,所述波束形成器被包括在所述第一助听器中,并且其中,由所述第一麦克风提供的所述电子输入信号被馈送到所述第一助听器中的所述波束形成器,并且其中,由所述第二麦克风提供的所述电子输入信号被传输到所述第一助听器中的所述波束形成器。

19. 根据权利要求 18 所述的助听器系统,其中,波束形成仅在所述第一助听器中执行。

20. 根据权利要求 18 所述的助听器系统,其中,波束形成器被包括在所述第二助听器中,并且其中,由所述第二麦克风提供的所述电子输入信号被馈送到所述第二助听器中的所述波束形成器,并且其中,由所述第一麦克风提供的所述电子输入信号被传输到所述第二助听器的所述波束形成器。

21. 根据权利要求 18 所述的助听器系统,其中,所述混合器被包括在所述第一助听器中。

22. 根据权利要求 9 所述的助听器系统,其中,所述混合器被包括在所述第一助听器中,并且其中,所述第二助听器包括第二混合器,并且其中,所述用户操作的界面可操作连接到所述第一助听器中的所述混合器和所述第二助听器中的所述第二混合器两者。

23. 根据权利要求 1 或 2 所述的助听器系统,其中,所述助听器系统包括第一压缩器。

24. 根据权利要求 7 所述的助听器系统,其中,所述第一助听器包括第一压缩器,并且所述第二助听器包括第二压缩器。

25. 根据权利要求 1 或 2 所述的助听器系统,其中,所述第一音频信号和第二音频信号的所述混合包括所述第一音频信号和所述第二音频信号的相加。

26. 根据权利要求 7 所述的助听器系统,其中,所述助听器系统包括信道特定滤波器,并且其中,借助于所述滤波器来对产生的噪声信道进行滤波,从延迟的信号参考减去所述滤波器的输出,并且其中,所述滤波器被选择来最小化均方差。

## 助听器中的波束形成

### 技术领域

[0001] 本发明总体上涉及具有波束形成能力的助听器系统,并且具体地涉及自适应双耳波束形成。

### 背景技术

[0002] 现代助听器的最重要任务之一是在存在噪声的情况下提供在言语可懂度上的改善。为了这个目的,已经广泛使用波束形成,特别是自适应波束形成,以便抑制干扰噪声。传统上,向助听器的用户提供在助听器中定向和全向模式之间改变的可能(例如,用户仅通过下述方式来改变处理模式:根据在特定环境中遇到的收听状况来在助听器上拨动扳钮开关或按动按钮以将所述装置置于优选模式中)。近来,已经在助听器中采用了用于在定向和全向模式之间切换的自动切换过程。

[0003] 依赖于特定的收听情况,全向和定向处理两者提供了相对于另一种模式的益处。对于较为安静的收听情况,全向处理通常优选于定向模式。这是因为,在所存在的任何背景噪声在振幅上相当低的情况下,全向模式应当提供对于周围环境中的各种声音的更多的访问,这可以提供对于环境的“连通性”,即被连接到外部世界的更大感觉。可预测当信号源在收听者侧面或后面时全向处理的一般偏好。通过声源对于收听者当前未面向的提供更多的访问,全向处理将改善对于从这些位置(例如,在服务员从收听者后面或侧面说话的餐馆中)到达的语音信号的识别。从除了收听者前面之外的位置到达的目标信号的全向处理的益处,在安静和噪声收听情况中都存在。对于收听者面向信号源(例如,感兴趣的说话者)的噪声收听状况,通过对于来自前面的信号的定向处理提供的提高的信噪比(SNR)有可能使得定向处理成为优选。在听力受损的收听者的每日体验中频繁地出现刚才所述的收听状况的每一个(在面向或未面向说话者的助听器用户的安静、噪声收听环境中)。因此,助听器用户定期地遇到定向处理将优选于全向模式的收听情况,并且反之亦然。

[0004] 在助听器的全向模式和定向模式之间的手动切换的方法的问题是,如果收听者不积极地切换模式,则收听者可能不知道模式的改变在给定的收听情况中会是有益的。另外,最适合的处理模式在一些收听环境中可以相当频繁地改变,并且收听者可能不能方便地手动切换模式以处理这样的动态收听状况。最后,许多收听者可能发现两种模式的手动切换和积极比较是烦人和不方便的。结果,他们可以将他们的装置永久地置于默认的全向模式中。

[0005] 然而,通过声音的无损编码来执行是由收听者手动地还是由听力仪器自动地选择定向麦克风。基本上,定向处理由其中增强一个声源(通常从0度起)并且衰减所有其他声源的空间滤波构成。因此,破坏了空间提示(spatial cue)。一旦去除了这个信息,则助听器或收听者不再能获得或检索到该信息。因此,在定向和全向模式之间手动或自动切换的这样的方法的主要问题之一是消除了对于收听者可能是重要的信息,这发生在听力仪器被切换到定向模式时。

[0006] 虽然定向模式的目的是提供感兴趣的信号的较好的信噪比,但是关于什么是感兴

趣的信号的判定最终是收听者的选择,而不能由听力仪器来判定。因为假定感兴趣的信号出现在收听者的观看方向上,则在收听者的观看方向外部出现的任何信号可以并且将通过定向处理来消除。这符合临床经验,临床经验表明当前市场上的自动切换算法未获得广泛的接受。病人一般更喜欢手动地切换模式,而不是依赖于这些算法的判定的切换模式。

## 发明内容

[0007] 因此,本发明的目的是提供一种助听系统,通过该助听系统,可以向用户同时提供定向和全向模式两者的益处。

[0008] 根据本发明,通过本发明的第一方面来实现上述和其他目的,本发明的第一方面涉及一种助听器系统,包括:第一麦克风和第二麦克风,用于提供电子输入信号;波束形成器,用于至少部分地基于所述电子输入信号来提供具有第一定向空间特性的第一音频信号(波束),其中,所述波束形成器进一步被配置成至少部分地基于所述电子输入信号来提供第二音频信号,所述第二音频信号具有与所述第一音频信号不同的另一个空间特性,并且其中,所述助听器系统进一步包括混合器,所述混合器被配置成用于混合所述第一音频信号和第二音频信号,以便提供要由用户听到的输出信号。

[0009] 通过将所述定向音频信号与具有另一个空间特性的音频信号混合以便提供要由用户听到的混合输出信号,用户获得了定向处理的益处(例如,感兴趣的信号的更好的可懂度),同时能够听到来自其他方向(多个)的声音。取决于混合比,即所述第二音频信号中的多少与所述第一音频信号混合,并且取决于所述第二音频信号的所述空间特性,用户将被提供有具有定向处理的益处的输出信号,并且同时用户感受到与周围声音联系更多的环境。

[0010] 根据优选实施例的所述助听器系统进一步包括处理器,所述处理器被配置成根据听力损害校正算法来处理所述混合信号。由此,保证所述混合信号具有将被用户听到的水平和频率特性。优选地,在所述助听系统中使用诸如扬声器的输出换能器(也称为接收器),以便将所述混合的音频信号转换为声音信号。

[0011] 根据本发明的第一方面的所述助听器系统可以替代地进一步包括处理器,所述处理器被配置成在混合所述第一音频信号和第二音频信号之前,根据听力损害校正算法来处理所述第一音频信号。因为通常用户主要对具有所述定向特性的所述第一音频信号感兴趣,所以这个替代实施例实现了根据所述用户的听力损害来至少处理用户最感兴趣的音频信号。

[0012] 根据本发明的一个实施例,所述波束形成器可以具有一个优选方向。例如,所述波束形成器可以具有由所述助听器系统的用户的“正前方”方向限定的一个优选方向,即,根据本发明的一个实施例,所述第一音频信号的定向特性可以具有被预定义在所述“正前方”方向中的方向。因此,限定了在“正前方”方向中的波束。根据替代实施例,在将波束方向保持固定的同时,所述第一音频信号的波束“宽度”或空间定向特性的形状可以是可适配的或至少可调整的。

[0013] 所述波束形成器优选地可以是自适应的,即,所述波束形成器根据具体情况来优化信噪比。

[0014] 通过使用可适配的波束形成器,实现很灵活的解决方案,其中,可以在用户移动所

述助听器系统的同时,聚焦在移动的声源上或聚焦在非移动的声源上。此外,可以更好地处理环境噪声状况的改变(例如,新的声源的出现、噪声源的消失或噪声源相对于助听器系统的用户的移动)。

[0015] 在根据本发明的第一方面的又一个优选实施例中,所述助听器系统可以包括用户操作的界面,所述用户操作的界面可操作地连接到所述混合器以用于控制所述第一音频信号和第二音频信号的混合。在此,实现了下述大优点:用户可以决定他/她可能想要听到周围声场的多少,并因此相对于他/她可能要感觉与周围的“联系”程度来上下调节。例如,如果本发明的助听器系统的用户处于宴会的情形,其中,他/她正在与面对他/她坐着的人进行会话,同时在多个其他参与者正在彼此交谈,则所述用户将位于声音环境中,所述声音环境经常被称为多谈话者乱哄哄噪声或仅称为乱哄哄噪声。在这样的情况下,本发明的助听器系统的用户将清楚地受益于定向处理,但是可能感到置身于在宴会的那组人的其余人之外,但是通过使用用于在第二音频信号的一些中混合的界面,将使得用户能够听到他/她可能选择的正在进行的其他会话,同时受益于相对于用户当前正在与其会话的人的定向处理。

[0016] 作为对于用户受控的替代或补充,可以独立于周围声音环境的分类来执行第一音频信号和第二音频信号的混合。这具有下述优点:可以优化在助听器系统中的音频信号处理,以处理特定的声音或噪声环境。

[0017] 优选地,可以将用户操作的界面置于独立的遥控设备中,所述独立的遥控设备例如类似于用于控制电视机的遥控设备,其经由无线链路可操作地连接到所述混合器。

[0018] 替代地,用户操作的界面可以包括手动可操作开关,所述手动可操作开关可以被置于所述助听器系统的壳体结构之中或之上。所述开关可以是扳钮开关或与在本领域中已知的助听器的音量滚轮相似的开关。替代地,所述开关可以被体现为接近传感器,所述接近传感器能够记录在所述传感器附近的手或手指移动。这样的接近传感器可以被体现为例如电容传感器。在另一个替代实施例中,所述开关可以是磁开关,诸如簧片开关、磁阻开关、巨型磁阻开关、各向异性磁阻开关或各向异性巨型磁阻开关。

[0019] 虽然许多听力受损的人遭受两耳的听力损失并且因此实际上使用两个助听器,但是双耳助听器系统的大多数在每一个助听器中独立地处理数据,而不交换信息。然而,近些年来,已经在助听器之间引入了无线通信,以便可以从一个助听器向另一个传送数据。因此,根据本发明的优选实施例,所述助听器系统可以是双耳助听器系统,所述双耳助听器系统包括经由通信链路而彼此互连的第一助听器和第二助听器,并且其中,所述第一麦克风位于所述第一助听器中,并且所述第二麦克风位于所述第二助听器中。由此实现了一种促进双耳波束形成的助听器系统。这除了别的之外进一步具有提高波束形成器的空间分辨率的优点,因为在耳朵之中或之处佩带第一助听器和第二助听器的一般成人的耳朵之间的距离大致在可听范围中的声音波长的数量级上。这将因此使得可以区分在空间上接近地定位的声源。然而,除了这些优点之外,双耳波束形成的一个关心问题是波束形成器仅产生一个信号,该信号有效地破坏了所有的双耳提示,诸如噪声的耳间时间差(ITD)和耳间水平差(ILD)。这些双耳提示对于使得人能够定位声源和/或区分声源是必要的。然而,通过混合第一音频信号和第二音频信号,可以保留双耳提示,同时向用户提供定向处理的益处。仿真已经显示,在根据本发明的助听器系统中,很大程度地保留了这些双耳提示(例如,参见关

于仿真结果的部分)。双耳助听器系统或用户可以确定给定情况期望的混合水平或混合比。

[0020] 根据双耳助听器系统的优选实施例,第一助听器 and 第二助听器中的每一个包括连接到所述波束形成器的另一个麦克风。由此,实现了将能够一次处理几个噪声源并且因此实现更好的噪声抑制的双耳助听器系统。

[0021] 根据双耳助听器系统的优选实施例,提供了一种手动可操作开关,用于控制第一音频信号和第二音频信号的混合,所述开关可以被置于所述第一助听器和 / 或第二助听器中,例如置于所述第一助听器和 / 或第二助听器的壳体结构中。

[0022] 根据又一个优选实施例,根据本专利说明书的描述的助听器系统可以是形成双耳助听器系统的一部分的单个助听器。

[0023] 根据优选实施例,由所述波束形成器产生的所述第一音频信号和第二音频信号的空间特性可以是基本上互补的。然而,虽然基本上互补,但是它们也可以在特定程度上重叠。这个实施例的一个大优点是,当将所述第二音频信号的提高部分与所述第一音频信号混合时,混合信号将从基本上定向的音频信号变为基本上全向的音频信号。因此,根据混合比,系统或用户可以在基本上定向和基本上全向的处理之间进行转换(例如,软切换),并且因此根据在任何给定情况下所期望的,受益于两者。

[0024] 替代地,第二音频信号的空间特性可以是基本上全向的。由此,实现了在计算上实现起来简单的系统,因为所述波束形成器仅提供具有定向空间特性的一个音频信号。

[0025] 根据替代优选实施例,所述第一音频信号和第二音频信号的空间特性以如下这样的方式由(所述波束形成器)产生:优选地当使用例如  $\beta = 1$  (下面在附图的详细描述下描述)的混合比的适当选择的混合比时,即当使用相等的加权来混合第一音频信号和第二音频信号时,结果产生的混合音频信号的空间特性是基本上全向的。

[0026] 可以根据用户的第一耳朵和 / 或第二耳朵的听力损失或根据周围声音环境的类别来执行所述混合本身。

[0027] 根据本发明,通过与一种助听器相关的本发明的第二方面来实现上述和其他目的,所述助听器包括:麦克风,用于提供定向音频信号和全向音频信号;处理器,其可操作地连接到所述麦克风,并且被配置用于提供要由用户听到的听力损害校正的输出信号,其中,所述助听器进一步包括混合器,用于混合所述定向音频信号和所述全向音频信号,由此提供混合的音频信号。

[0028] 根据本发明的第二方面的实施例进一步涉及一种助听器,所述助听器包括用户操作的界面,所述用户操作的界面可操作地连接到所述混合器,由此所述混合可以被用户控制。

[0029] 根据本发明的第二方面的实施例,所述听力损害校正的输出信号可以基于混合音频信号或定向音频信号或全向音频信号。

[0030] 根据本发明的第二方面的实施例的一种助听器可以被配置用于形成双耳助听器系统的一部分。

[0031] 根据本发明,通过与一种双耳助听器系统相关的本发明的第三方面来实现上述和其他目的,所述双耳助听器系统包括:第一助听器,其具有用于提供定向音频信号的定向麦克风系统和用于提供第一听力损害校正的输出信号的处理器;第二助听器,其具有用于提供全向音频信号的全向麦克风系统和用于提供第二听力损害校正的输出信号的接收器,其

中,经由在所述第一助听器和第二助听器之间的双向通信链路,所述第一助听器适于接收基于所述全向音频信号的音频信号,并且所述第二助听器适于接收基于所述定向音频信号的音频信号,其中,所述第一助听器进一步包括第一混合器,第一混合器用于混合基于所述全向和所述定向音频信号的信号,以便提供第一混合信号,并且其中,所述第二助听器进一步包括第二混合器,第二混合器用于混合基于所述全向和所述定向音频信号的信号,以便提供第二混合信号。

[0032] 在根据本发明的第三方面的实施例中,由所述第一混合器和/或第二混合器执行的所述混合可以基于从所述全向麦克风系统和/或所述定向麦克风系统得出的信号的类别。

[0033] 在根据本发明的第三方面的另一个实施例中,可以根据从所述全向麦克风系统和/或所述定向麦克风系统得出的信号的目标信噪比(SNR)和/或信号压力级(SPL)来执行所述混合。

[0034] 根据本发明的第三方面的所述双耳助听器系统可以进一步包括用户操作的界面,所述用户操作的界面可操作地连接到所述第一混合器和/或第二混合器。

[0035] 依照根据本发明的第三方面的双耳助听器系统的又一个实施例,所述第一听力损害校正的输出信号可以至少部分地基于所述第一混合信号。作为其补充或替代地,所述第二听力损害校正的输出信号可以至少部分地基于所述第二混合信号。

[0036] 根据本发明的第三方面的实施例的第一混合信号和第二混合信号基本上相同,或者,可以根据相同的混合比来执行所述混合。

[0037] 在根据本发明的第三方面的优选实施例中,可以根据与用户的第一耳朵相关联的听力损失来产生所述第一听力损害校正的输出信号,并且,可以根据与用户的第二耳朵相关联的听力损失来产生所述第二听力损害校正的输出信号。

[0038] 根据本发明的第二或第三方面的实施例,可以根据用户的第一耳朵和/或第二耳朵的听力损失来执行所述混合。

[0039] 虽然以上已经描述了本发明的三个方面的几个实施例,但是应当明白,来自该三个方面之一的实施例的任何特征可以被包括在其他两个方面之一或两者的实施例中,并且当在本专利申请中其被称为“实施例”时,可以明白其可以是根据本发明的三个方面的任何一个的实施例。

## 附图说明

[0040] 下面,将参考附图更详细地解释本发明的优选实施例,其中:

[0041] 图1示出根据本发明的一方面的助听器系统的实施例;

[0042] 图2示出根据本发明的一方面的助听器系统的替代实施例;

[0043] 图3示出根据本发明的一方面的助听器系统的又一个替代实施例;

[0044] 图4示出根据本发明的一方面的双耳助听器系统;

[0045] 图5示出根据本发明的一方面的双耳助听器系统的替代实施例;

[0046] 图6图示对于在图4中所示的双耳助听器系统的替代实施例;

[0047] 图7图示对于在图5中所示的双耳助听器系统的替代实施例;

[0048] 图8图示具有定向空间特性的第一音频信号与具有与第一音频信号的空间特性

不同的空间特性的另一个音频信号的混合；

[0049] 图 9 图示在仿真中的根据本发明的一些方面的助听器系统的频率相关性能；

[0050] 图 10 图示在仿真中的根据本发明的一些方面的助听器系统的角度相关性能；

[0051] 图 11 图示作为入射角的函数的分别用于单个和多个噪声源的耳间时间差上的误差；以及

[0052] 图 12 图示作为入射角的函数的估计的耳间水平差。

## 具体实施方式

[0053] 下文中将参考示出了本发明的示例性实施例的附图更全面地描述本发明。然而，本发明可以以不同形式来体现，并且不应当被解释为限于在此阐述的实施例。而是，这些实施例被提供使得本公开将是透彻和完整的，并且将向本领域内的技术人员全面地传达本发明的范围。在全部附图中，相同的附图标记指的是相同的元件。因此，将不参考每一个附图的描述来详细地描述相同的元件。

[0054] 图 1 示出根据本发明的一方面的助听器系统的实施例。所图示的助听器系统被体现为助听器 2，助听器 2 包括两个麦克风 4 和 6，分别用于提供电子输入信号 8 和 10。所图示的助听器 2 也包括波束形成器 12，波束形成器 12 被配置用于提供具有定向空间特性的第一音频信号 14（有时称为波束）。第一音频信号 14 至少部分地基于电子输入信号 8 和 10，并且，第二音频信号 16 也至少部分地基于电子输入信号 8 和 10。波束形成器 12 也被配置用于提供第二音频信号 16，第二音频信号 16 具有与第一音频信号 14 的空间特性不同的空间特性。所述第一音频信号和第二音频信号 14 和 16 在混合器 18 中混合，以便提供混合的音频信号 20。助听器 2 进一步包括压缩器 22，压缩器 22 被配置用于根据听力损害校正算法来处理混合的音频信号 20。听力损害校正的混合信号随后被所图示的接收器 24 变换为声音信号。优选地，在诸如数字信号处理器 (DSP) 26 的信号处理器中包括波束形成器 12、混合器 18 和压缩器 22。可以明白，可以以软件来实现下述单元的任何一个或全部：波束形成器 12、混合器 18 和压缩器 22。此外，可以以软件来实现单元 12、18 和 22 的一些部分，而以诸如 ASIC 的硬件来实现其他部分。因为大多数听力障碍是频率相关的，所以优选地，压缩器 22 可以被配置成根据听力损害校正算法来执行混合的音频信号 20 的频率相关的处理。优选地，根据助听器 2 的用户的特定的估计或测量的听力损害来选择或产生这个听力损害校正算法。

[0055] 在图 1 中也示出了(可选的)用户操作的界面 28，用户操作的界面 28 经由控制链路 30 可操作地连接到混合器 18。在一个实施例中，所图示的用户操作的界面 28 可以包括在助听器 2 的壳体结构(未示出)上的诸如音量滚轮的致动器或传感器(未示出)。这将因此使得用户能够通过下述方式来控制第一音频信号和第二音频信号 14 和 16 的混合：使用他/她的手或手指来手动地激活致动器或传感器(未示出)。在另一个实施例中，所图示的用户操作的界面 28 形成遥控设备的一部分，无线控制信号 30 可以从这个遥控设备被发送到助听器 2，并且在助听器 2 处被接收，以便控制在混合器 18 的第一音频信号和第二音频信号 14 和 16 的混合。在这个实施例中，可以明白，助听器 2 配备有用于从遥控设备接收无线控制信号的装置，虽然在图 1 中未明确地示出这些特征。

[0056] 此外，可以明白，所图示的助听器 2 可以是耳后类型的助听器、耳内类型的助听

器、完全在耳道内类型的助听器或接收器在耳朵内类型的助听器(即,一种类型的助听器,其中,除了接收器 24 之外的在图 1 中所示的所有特征被置于壳体机构中,该壳体结构被配置用于被置于用户耳后,并且其中,接收器 24 被置于耳塞中,所述耳塞例如可以是耳模,所述耳塞被配置用于被置于用户的耳道或耳甲腔中)。

[0057] 图 2 示出在图 1 中所示的根据本发明的一方面的助听器系统的替代实施例。在图 1 和 2 中所示的实施例之间的唯一差别是分类器 32。通过包括分类器 32,可以让助听器 2 执行第一音频信号和第二音频信号 14 和 16 的自动混合,其中,可以对于不同的收听情况优化该混合。例如,如果除了可能用户感兴趣的一个声源之外周围声音环境是安静的,则可以以结果产生的混合的音频信号 20 基本上是全向的这样的方式来执行混合。

[0058] 然而,因为不可能先验地考虑所有可能的收听情况,因此不可能优化在任何可能的收听情况下对于用户最佳的混合,所以用户可以否决由分类器 32 控制的自动混合。用户可以通过激活用户操作的界面 28 来如此进行。

[0059] 在图 2 中所示的助听器 2 的更简化的实施例中,仅根据分类器 32 对于周围的声音环境的分类来执行混合。因此,这样的实施例不包括用户操作的界面 28。因此,在这个简化实施例中,用户不能否决由分类器 32 控制的混合。

[0060] 图 3 示出了根据本发明的一方面的助听器系统的替代实施例。所图示的助听器系统被体现为助听器 2,并且在许多方式上类似于图 1 和 2 中图示的实施例。因此,仅将详细描述这些实施例的差别。在所图示的实施例中,压缩器 22 被配置用于根据听力损害校正算法来处理第一音频信号 14,以便提供听力损害校正输出信号 34。这在某些情况下是有利的,因为波束形成的音频信号 14 通常被引导至用户感兴趣的声源。因此,用户将有兴趣听到对于他/她方便的大声和清楚的特定声源。然而,为了使得用户有可能也从其他方向听到声音并且因此感到联系到周围的声音环境,将信号 34 与第二音频信号 16 混合,以便提供混合输出信号 36,混合输出信号 36 在接收器 24 中被转换为声音。如图所示,助听器系统也可以包括(可选的)用户操作的界面 28,通过用户操作的界面 28,用户可以以与如上所述类似的方式控制混合。

[0061] 在本发明的替代实施例中,在图 1-3 的任何一个中图示的助听器 2 可以包括一个或两个附加麦克风,使得它总共可以包括 3 或 4 个麦克风或甚至比 4 个更多的麦克风。

[0062] 在另一个实施例中,参考在图 1-3 中所示的实施例的任何一个所述的助听器 2 可以被配置用于形成双耳助听器系统的一部分,该双耳助听器系统包括另一个助听器。进一步可以彼此协调在形成双耳助听器系统的一部分的两个助听器中的信号处理。

[0063] 图 4 示出了根据本发明的另一个实施例的助听器系统,其中,所述助听器系统是双耳助听器系统,包括:第一助听器 2,其具有一个麦克风 4;以及,第二助听器 38,其包括第二麦克风 6。第二助听器 38 进一步包括压缩器 40 和接收器 42。在所图示的双耳助听器系统中,仅在助听器 2 中执行波束形成。因此,由第二助听器 38 提供的电子输入信号 10 被传输到第一助听器 2 中的波束形成器 12,如虚线箭头 44 所指示的。以与以上参考在图 1-3 中所示的实施例解释的方式类似的方式来执行助听器 2 中的电子输入信号 8 和 10 的进一步的处理,包括音频信号 14 和 16 的混合。然而,重要的差别是,混合的音频信号 20 也被传输到第二助听器 38 的压缩器 40,如虚线箭头 46 所指示的。优选地,压缩器 40 根据听力损害校正算法来处理混合的音频信号,以便补偿用户的第二耳朵的听力损害。然后,来自压缩器

40 的输出信号被馈送到第二接收器 42, 第二接收器 42 被配置用于将压缩器的输出信号变换为要被用户听到的声音信号。因为遭受听力障碍的许多人两耳都遭受听力损失, 并且在许多情况下甚至两耳遭受不同的听力损失, 所以优选地, 压缩器 22 被配置用于根据听力损害校正算法来处理混合的音频信号 20, 以便减轻用户的第一耳朵的听力损失, 而第二助听器 38 的压缩器 40 被配置用于根据听力损害校正算法来处理混合的音频信号 20, 以便减轻用户的第二耳朵的听力损失。

[0064] 虽然未明确地图示, 但是输入信号 10 可以在助听器 38 中进行附加信号处理。

[0065] 如本领域中已知的, 通过有线或无线链路 (例如, 双向链路) 来促进两个助听器 2 和 38 之间的、如虚线箭头 44 和 46 所指示的信号 10 和 20 的传输。

[0066] 图 5 示出根据本发明的一方面的替代助听器系统, 该替代助听器系统在此被体现为双耳助听器系统, 其包括第一助听器 2 和第二助听器 38。所图示的助听器 2、38 中的每一个包括: 麦克风 4、6; 波束形成器 12、48; 混合器 18、50; 压缩器和接收器 24、42。在助听器 2 中, 波束形成器 12、混合器 18 和压缩器 22 形成了诸如数字信号处理器 (DSP) 26 的信号处理单元的一部分。对应地, 在助听器 38 中, 波束形成器 48、混合器 50 和压缩器 40 形成了诸如数字信号处理器 (DSP) 54 的信号处理单元的一部分。

[0067] 第一助听器 2 的麦克风 4 提供电子输入信号 8, 该电子输入信号 8 被馈送到波束形成器 12, 并且也被传输到第二助听器 38 的波束形成器 48, 如虚线箭头 60 所指示的。类似地, 第二助听器 38 的麦克风 6 提供电子输入信号 10, 该电子输入信号 10 被馈送到波束形成器 48 并且也被传输到第一助听器 2 的波束形成器 12, 如虚线箭头 62 所指示的。因此, 波束形成器 12 和 48 中的每一个接收由两个麦克风提供的电子信号。以以上相对于在图 1-3 中所示的实施例描述的方式类似的方式来执行在助听器 2、38 的每一个中的电子输入信号 8、10 的进一步处理。可以通过例如双向有线或无线链路来促进在助听器 2、38 之间的输入信号 8、10 的传送, 如虚线箭头 60、62 所指示的。

[0068] 在图 5 中图示的双耳助听器系统的一个实施例中, 第一助听器和第二助听器 2、38 的波束形成器 12、48 可以被配置成以如下这样的方式来执行协调的波束形成: 音频信号 14 和 56 基本上相同和 / 或音频信号 16 和 58 基本上相同。实现向在两个助听器中的混合器 18、50 的输入信号的方式是类似的。如参考图 4 解释的, 压缩器 22 和 40 被配置成分别根据用户的第一耳朵和第二耳朵的听力损失来处理混合的音频信号 20 和 64。

[0069] 在图 5 中也示出了 (可选的) 用户操作的界面 28。所图示的用户操作的界面 28 可操作地连接到第一助听器 2 中的混合器 18, 如虚线箭头 30 所指示的, 并且所图示的用户操作的界面 28 可操作地连接到第二助听器 38 中的混合器 50, 如虚线箭头 52 所指示的。在优选实施例中, 用户操作的界面 28 形成遥控设备的一部分, 由此, 可以通过无线链路来促进用户操作的界面 28 与助听器 2 和 38 之间的操作连接, 通过该无线链路, 控制信号可以被发送到两个助听器 2 和 38 中的每一个。在优选实施例中, 用户可以通过适当地激活用户操作的界面 28 来独立于彼此地控制两个助听器 2 和 38 的每一个中的混合。在另一个实施例中, 用户操作的界面 28 被配置用于在两个助听器 2 和 38 的每一个中提供协调和类似数量的混合。在又一个替代实施例中, 用户操作的界面 28 被包括在置于助听器 2 和 38 之一或两者的壳体结构 (未示出) 中的开关结构中。所述开关结构可以例如包括机械致动器或接近传感器或在发明内容中描述的任何其他类型的开关结构。在另一个实施例中, 用户操作的

界面 28 可以由两个独立部分构成,一个独立部分用于控制助听器 2 中的混合,并且一个独立部分用于控制助听器 38 中的混合。在此,可以明白,用户操作的界面 28 也可以包括开关结构(未示出)的两个独立部分,其中每一个可以被置于两个助听器 2 或 38 的每一个中。因此,以这种方式,可以通过助听器 2 中的开关(未示出)来控制助听器 2 中的混合,并且可以通过助听器 38 中的开关(未示出)来控制助听器 38 中的混合。

[0070] 图 6 图示了与在图 4 中所示的双耳助听器系统类似的双耳助听器系统,但是现在其中,助听器 2、38 的每一个已经分别配备有一个附加麦克风 5 和 7。因此,将仅描述图 6 和图 4 中所示的实施例之间的差别:助听器 2 中的附加麦克风 5 提供电子输入信号 9,电子输入信号 9 被馈送到波束形成器 12,并且助听器 38 中的附加麦克风 7 提供电子输入信号 11,电子输入信号 11 经由有线或无线链路被传输到助听器 2 中的波束形成器 12,如虚线箭头 45 所指示的。在此,波束形成器 12 将具有要处理的四个麦克风信号,由此更准确和精确的波束形成是可能的(如下所述)。

[0071] 如本领域中已知的,可以通过有线或无线链路(例如,双向链路)来促进两个助听器 2 和 38 之间的、如虚线箭头 44、45 和 46 所指示的信号 10、11 和 20 的传输。

[0072] 类似地,图 7 图示了与图 5 中所示的双耳助听器系统类似的双耳助听器系统,但是现在其中,助听器 2、38 的每一个已经分别配备有一个附加麦克风 5 和 7。因此,将仅描述图 7 和图 5 中所示的实施例之间的差别:助听器 2 中的附加麦克风 5 提供电子输入信号 9,电子输入信号 9 被馈送到波束形成器 12 并且优选地经由有线或无线链路被传输到助听器 38,如虚线箭头 61 所指示的,其中,它(9)被馈送到助听器 38 中的波束形成器 48。类似地,助听器 38 中的附加麦克风 7 提供电子输入信号 11,电子输入信号 11 被馈送到波束形成器 48 并且经由链路(优选地,无线链路)被传输到助听器 2 中的波束形成器 12,如虚线箭头 63 所指示的。据此,波束形成器 12 和波束形成器 48 两者将具有要处理的四个麦克风信号,由此,更准确和精确的波束形成是可能的(如下所述)。此外可以彼此协调由两个波束形成器 12 和 48 执行的波束形成。

[0073] 可以通过例如双向有线或无线链路来促进助听器 2、38 之间的输入信号 8、9、10 和 11 的传输,如虚线箭头 60、61、62 和 63 所指示的。

[0074] 可以明白,优选地,图 1-7 的任何一个中所示的波束形成器 12、48 是自适应的。此外,可以明白,图 3-7 的任何一个中图示的助听器 2、38 的每一个可以包括如参考图 2 所述的分类器(未示出)。

[0075] 图 8A-8C 图示了具有定向空间特性 66 的第一音频信号与具有与第一音频信号的空间特性 66 不同的空间特性 68 的另一个音频信号的混合,以便提供混合信号。

[0076] 图 8A-8C 中图示的空间特性被给出为极坐标图,该极坐标图示出了在基本上水平的平面中的作为角度的函数的周围声场的放大。图 8A 中图示的混合示出了用户感兴趣的谈话者被置于 0 度角的情况,并且干扰噪声源被置于 90 度角。空间特性 66 是由波束形成器提供的语音估计,并且空间特性 68 是由波束形成器提供的噪声估计。图 8A 中图示的空间特性的最后一列示出了对于因子  $\beta$  的各个值而言的结果产生的混合信号的空间特性(例如参见下面的等式(16)来获得更多的细节)。因子  $\beta$  图示噪声估计中的多少与语音估计混合。因此, $\beta = 1$  的值对应于所有的噪声估计与语音估计混合的情况,产生全向混合信号,并且另一种极端情况,其中, $\beta = 0$  的值对应于没有噪声估计与语音估计混合的情况,因此

产生具有等于语音估计的空间特性的混合信号。在图 8A 的最后一列中也图示了两种中间情况,两种中间情况示出了  $\beta = 0.3$  和  $\beta = 0.7$  的混合信号的空间特性。在本发明的优选实施例中,用户能够控制混合因子  $\beta$ ,使得他/她可以决定他/她可能想要听到多少噪声估计,由此控制对于周围声音环境的“连通性”。

[0077] 在图 8B 和 8C 中图示了与以上参考图 8A 描述的情况类似的情况,但是具有下述差别:在图 8B 中,将干扰噪声源置于 110 度角,并且在图 8C 中,将干扰噪声源置于 180 度角。

[0078] 图 8A-8C 的任何一个中图示的混合仅示出了可以由图 1-7 的任何一个图示的混合单元 18 或 50 执行的混合的两个简单示例。可以想象除了如图 8A-8C 中图示的仅仅相加之外的其他种类的混合,例如一些适当的加权和相乘,并且,展现不同的空间特性的其他音频信号的混合也是可能的。因此,根据所使用的混合比(即第一信号和第二信号如何相对于彼此加权)和所产生的第一音频信号和第二音频信号的空间特性,可以获得混合信号的任何期望的空间特性。

[0079] 下面,将以数学方式来描述由图 1-7 的任何一个中图示的波束形成器 12 和/或 48 的任何一个执行的波束形成的方法的示例:

[0080] 考虑由下式描述的在时间  $t$  的入射声场:

$$[0081] \quad y(r, t) = s(t - \alpha \cdot r) + w(r, t) \quad (1)$$

[0082] 其中,  $s(t)$  是具有缓慢度  $\alpha$  (根据本发明的优选实施例,缓慢度被定义为由在中间的声速划分的传播方向)的感兴趣的传播平面波(即,表示用户感兴趣的信号),并且其中,  $w(r, t)$  表示干扰噪声场。在场的自变量中包括  $r$  和  $t$  表示它们依赖于空间和时间。在  $M$  个空间位置(对应于  $M$  个空间麦克风位置)采样入射波场,因此产生  $M$  个时间信号

$$[0083] \quad y_m(t) = s(t - \alpha \cdot r_m) + w(r_m, t) \quad (2)$$

[0084] 波束形成器然后对齐所测量的响应,使得感兴趣的信号同相

$$[0085] \quad z_m(t) = y_m(t + \alpha \cdot r_m) = s(t) + w_m(t) \quad (3)$$

[0086] 其中,  $w_m(t) = w(r_m, t + \alpha \cdot r_m)$ 。可以将对应的采样信号模型写为:

$$[0087] \quad z_m(n) = s(n) + w_m(n) \quad (4)$$

[0088] 然后产生  $M-1$  个噪声信道

$$[0089] \quad v_m(n) = z_0(n) - z_x(n), m \neq 0 \quad (5)$$

[0090] 噪声信道被以向量形式写入,并使用具有  $N$  个抽头的信道特定滤波器来滤波,并且从延迟的信号参考(第一信道)减去输出

$$[0091] \quad e(n) = z_0(n - N/2) - \sum_{m=1}^{M-1} \mathbf{h}_m^T \mathbf{v}_m(n) \quad (6)$$

[0092] 其中,  $(\cdot)^T$  是  $(\cdot)$  的转置,并且

$$[0093] \quad \mathbf{h}_m = (h_m(0) \cdots h_m(N-1))^T \quad (7)$$

$$[0094] \quad \mathbf{v}_m(n) = (v_m(0) \cdots v_m(n-N+1))^T \quad (8)$$

[0095] 等式(6)可以被更紧凑地写为

$$[0096] \quad e(n) = z_0(n - N/2) - \mathbf{h}^T \mathbf{v}(n) \quad (9)$$

[0097] 其中

$$[0098] \quad \mathbf{h} = (\mathbf{h}_1^T \cdots \mathbf{h}_{M-1}^T)^T \quad (10)$$

$$[0099] \quad \mathbf{v}(n) = (\mathbf{v}_1^T(n) \dots \mathbf{v}_{M-1}^T(n))^T \quad (11)$$

[0100] 滤波器被选择来最小化均方差

$$[0101] \quad h_{\text{opt}} = E\{e(n)^2\} \quad (12)$$

[0102] 可以明白,可以使用作为 LMS (最小均方) 的更新方案来在线完成这一点,或可以在适当的情况下计算滤波器,并且对于特定的噪声情况而言,滤波器可以是固定的。

[0103] 假定感兴趣的信号与噪声无关(这在大多数情况下有意义,因为感兴趣的信号通常是与干扰噪声无关的语音信号),以选择滤波器的这种方式来产生噪声处理的估计  $w_0(n)$  :

$$[0104] \quad \hat{w}_0(n - N/2) = \mathbf{h}^T \mathbf{v}(n) \quad (13)$$

[0105] 并且,从这个结果接下来是

$$[0106] \quad \hat{s}(n) = z_0(n) - \hat{w}_0(n) \quad (14)$$

[0107] 并且

$$[0108] \quad \hat{w}_m(n) = \hat{w}_0(n) - v_m(n), m \neq 0 \quad (15)$$

[0109] 如果假定可以以足够的精度来估计噪声处理  $w_0(n)$ , 则如在(14)和(15)中所示,也可以提取其他四个信号。

[0110] 现在可以通过下式来找到独立的信道的修正估计

$$[0111] \quad x_m(n) = \hat{s}(n) + \beta_m \hat{w}_m(n) \quad (16)$$

[0112] 其中,  $\beta_m$  是用于控制不同信道的信噪比(即噪声估计中的多少与语音估计混合)的参数。

[0113] 仿真结果

[0114] 已经在仿真中测试了所述方法,其中,将根据本发明的一方面的双耳助听器系统(以下称为双耳波束形成器)与根据本发明的另一个方面的未处理的信号和单耳自适应波束形成器作比较。在仿真中,使用自由场模型,并且假定远场传播,即,声学模型基于远场近似。该阵列具有四个麦克风,在头的任一侧有两个,即对应于根据本发明的一方面的双耳助听器系统,根据本发明的一方面的双耳助听器系统包括两个助听器,每一个配备有两个麦克风,即前麦克风和后麦克风。在单独的助听器上的麦克风之间的距离是 1cm,并且在两个前麦克风之间的距离是 14cm,而在两个后麦克风之间的距离是 15cm。假定声速是 342m/s,并且整个双耳助听器系统的采样频率是 16kHz。与特定的噪声信道  $h_m$  相关联的滤波器具有 21 个抽头,导致目标信号的 10 个采样的处理延迟。从 0 度播放语音信号。假定热噪声是具有高斯分布的在空间和时间上是白的。调整噪声水平,使得信噪比是 30dB (对应于 60dB 的声压级和 30dB 的麦克风噪声水平)。

[0115] 频率相关的性能:

[0116] 在这个仿真中,仅使用一个干扰源。干扰源在该情况下是带宽受限的定向噪声分量。与麦克风阵列作比较,入射角是 90 度。噪声分量的带宽是 1kHz,并且与来自前面的目标信号无关。噪声分量的中心频率从 500Hz-7.5kHz 改变。在该情况下,参数  $\beta$  被选择来给出噪声的最大衰减( $\beta_m = 0$ )。可以在图 9 中看到结果。曲线 78 描述了在(全向)麦克风中的任何一个上的未处理的信号,曲线 80 示出了单耳助听器的信噪比,并且曲线 82 是双耳

助听器系统的结果。对于低频,双耳助听器系统胜过单耳助听器,而对于高频而言,差异较小。

[0117] 角度相关性能:

[0118] 在这个仿真中也使用仅一个干扰源。在该情况下,干扰源是带宽受限的定向噪声分量。噪声的中心频率是 2kHz,并且噪声分量的带宽是 1kHz,并且与来自前面的目标信号无关。入射角从 0-90 度改变。在该情况下,参数  $\beta$  也被选择来给出噪声的最大衰减( $\beta_m = 0$ )。可以在图 10 中看到结果。曲线 84 描述了在麦克风中的任何一个上的未处理的信号,曲线 86 示出了单耳助听器的信噪比,并且曲线 88 是双耳助听器系统的结果。对于在 0 和 90 度之间的角度,双耳助听器具有比单耳助听器好得多的性能,而所述两个系统在后半球中显示类似的性能。

[0119] 多个噪声源

[0120] 具有更多麦克风的益处之一是波束形成器具有要使用来工作的更大的自由度。因此,执行另一个仿真,以便示出多个来源的性能差别。对于这个仿真,从 90、120 和 180 度入射三个干扰源。所有的噪声源的中心频率被选择为 2kHz,并且带宽是 1kHz。噪声源相互无关,并且与目标信号无关。在表 1 中,可以看到三种测试情况的信噪比。在此,以大约 29dB 的 SNR 增益而言,双耳助听器系统的优点是显然的,而单耳助听器仅给出了 8dB 的 SNR 增加。

[0121]

| 方法  | SNR     |
|-----|---------|
| 未处理 | -4. 8dB |
| 单耳  | 2. 5dB  |
| 双耳  | 24. 5dB |

[0122] 表 1

[0123] 弥散噪声的性能

[0124] 弥散噪声的性能对于助听器应用很有趣,因为在诸如会议室、餐饭馆或咖啡厅中的高度回响的环境中经常遇到这样的噪声场。因此,也执行对于弥散噪声的仿真,其中,将弥散噪声场仿真为:

$$[0125] \quad d(\mathbf{r}, t) = \sum_{i=0}^{I-1} g(t) * p(t - \mathbf{a}_i \cdot \mathbf{r}) \quad (17)$$

[0126] 其中,  $g(t)$  是与  $p(t)$  的延迟版本卷积的具有 6kHz 的截止频率的线性相位低通滤波器,  $p(t)$  是具有零平均值和高斯分布的白随机时间信号。通过下式来给出变量  $\mathbf{a}_i$

$$[0127] \quad \mathbf{a}_i = (\sin \theta_i \cos \theta_i)^T / c \quad (18)$$

[0128] 其中,  $\theta_i$  是在间隔  $[0, 2\pi]$  上具有均匀分布的随机入射角,并且  $c$  是声速。波的数目被选择为  $I = 2000$ 。弥散波场在麦克风的位置被评估,并且被采样以产生离散时间噪声序列。可以在表 2 中看到不同测试情况的结果。

[0129]

| 方法  | SNR    |
|-----|--------|
| 未处理 | -3.3dB |
| 单耳  | 0.57dB |
| 双耳  | 3.0dB  |

[0130] 表 2

[0131] 可以注意到,对于双耳和单耳助听器而言,性能增益比定向噪声情况小得多。单耳助听器的 SNR 增益是大约 4dB,并且对于双耳助听器系统是 6dB。

[0132] 重要的本地化提示是耳间时间差(ITD)和耳间水平差(ILD)。因此,已经通过仿真来研究了这些双耳提示。

[0133] 耳间时间差

[0134] 首先,通过仿真来研究再现定向噪声源的正确 ITD 的能力。在第一仿真中,在波场中存在单个噪声分量。噪声的中心频率被选择为 2kHz,并且噪声分量的带宽被选择为 1kHz,并且与来自前面的目标信号无关。入射角从 10-350 度变化。计算在右耳上的信道与在左耳上的对应信道之间的 ITD。通过找到两个不同通道的噪声估计的交叉相关函数中的内插峰值来实现这一点。将这个值与定向噪声分量的真实 ITD 作比较。在图 11 中的曲线 90 中示出了以微秒计的误差。由于所研究的两个麦克风的线性阵列几何形状而造成误差围绕 0 和 180 度是对称的。

[0135] 执行对应的仿真,其中,另外两个无关的干扰源也是活动的。噪声源从 90 和 180 度入射,并且具有与所研究的噪声源相同的空间特性。再一次,计算源的估计 ITD 和真实 ITD 之间的 ITD 误差。结果被显示为图 11 中的曲线 92。可以看出,与单个噪声源情况作比较,ITD 误差对于多个噪声情况而言更大。然而,与在 ms 量级上的耳间的真实 ITD 作比较,误差仍然很小。

[0136] 耳间水平差

[0137] 也相对于 ILD 测试波束形成方法。在波场中存在单个噪声分量。噪声的中心频率被选择为 2kHz,并且噪声分量的带宽是 1kHz,并且与来自前面的目标信号无关。入射角从 10-350 度变化。在组合语音信号和噪声信号之前,将在头部右侧上的噪声信号乘以因子 1/2。通过下述方式来估计 ILD:提取在头部两侧上的噪声分量,并且计算各自的自动相关函数的最大值的比率。在图 12 中,通过曲线 94 给出估计的 ILD,并且通过直线 96 来给出真实的 ILD。仿真显示,该波束形成方法能够再现波场的正确的 ILD。

[0138] 在本专利说明书中描述了一种用于助听器的自适应波束形成算法,其中,在头部的相对侧上的助听器之间具有双耳耦合。然而,应当明白,也可以使用非自适应波束形成算法。当设计双耳算法时的关键考虑问题之一是:虽然波束形成器应当抑制不想要的定向干扰,但是其不应当破坏由根据本发明的助听器系统的用户用来目标定位的干扰的双耳提示。

[0139] 所提出的算法产生对从目标方向入射的信号估计(通常被选择为在 0 度固定),但是也给出在所有的麦克风上的噪声分量的估计。在输出处提供的信号(其然后被传递以

在助听器进行进一步的处理)是目标信号和噪声的适当混合。混合比可以被用户通过遥控器调整,或由给定的当前声学环境下的助听器来决定。

[0140] 在本专利说明书中提供的仿真仅与定向噪声抑制性能相关,即仅与目标信号和没有噪声混合相关,并且与具有自适应波束形成的单个助听器作比较。当仅存在一个定向噪声源时,显示单耳助听器执行得比如果未应用波束形成更好,但是,对于所有角度,特别是在前半部,双耳助听器系统执行得比单耳助听器好得多。这适用于噪声的不同频率。在此,性能增益在低频率中是最大的。当在场中存在三个定向噪声源时,单耳助听器的性能增益是 8dB。这是在阵列中的麦克风的小数目(仅为 2)不能适当地抑制这个数目的来源的结果。然而,双耳阵列(具有 4 个麦克风)实现了 28dB 的 SNR 增益。对于弥散噪声场也执行了仿真。然而,波束形成算法的性能降低,分别地,单耳助听器的 SNR 增益是 4dB,并且双耳助听器系统的 SNR 增益是 6dB。

[0141] 也评估了所提出的用于再现干扰噪声的 ITD 和 ILD 的算法的能力。显示对于单个干扰源的情况和对于多个干扰噪声源的情况而言,所估计的 ITD 的误差处于微秒量级。必须将此认为较小,因为真实的 ITD 在微秒范围中。也显示,当单个干扰源在头部两侧上产生不同的压力级时,正确地再现 ILD。

[0142] 因此,如上所述,音频信号的波束形成和混合在助听器系统中的使用上是可行和有利的。然而,本领域内的技术人员将明白,在不偏离本发明的精神或必要特性的情况下,可以以除了如上所述并且在附图中图示的那些之外的其他具体形式来体现本发明,并且本发明可以利用多种不同算法中的任何一种。例如,算法的选择通常是应用特定的,该选择取决于多种因素,其中包括预期的处理复杂度和计算负荷。因此,在此的公开和描述意欲是在权利要求中阐述的本发明的范围的说明,而不是限定。

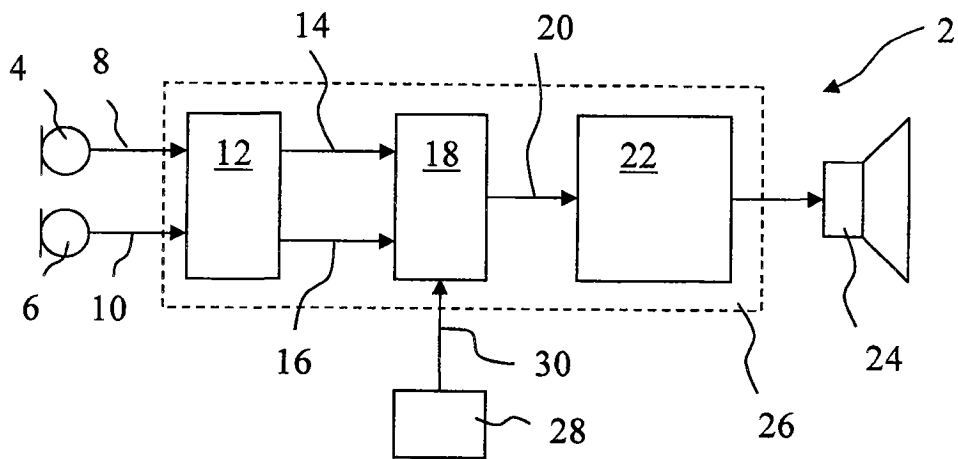


图 1

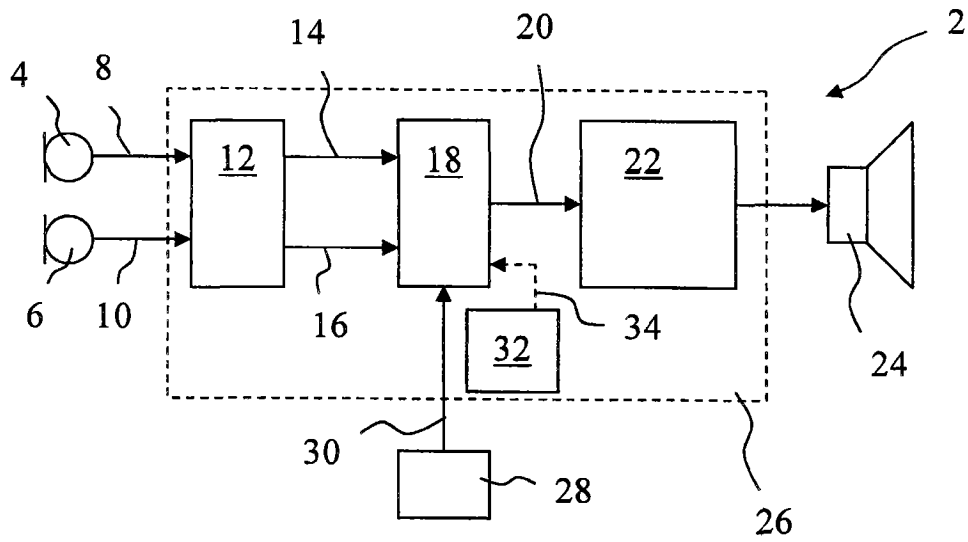


图 2

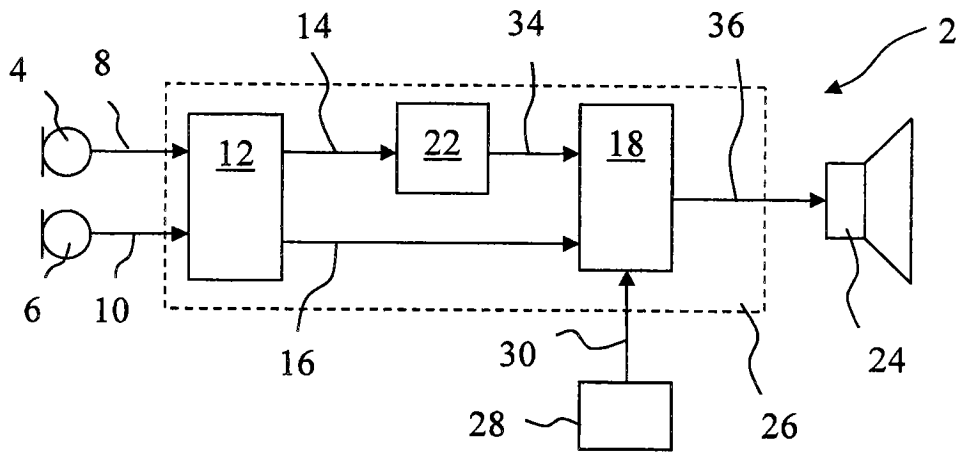


图 3

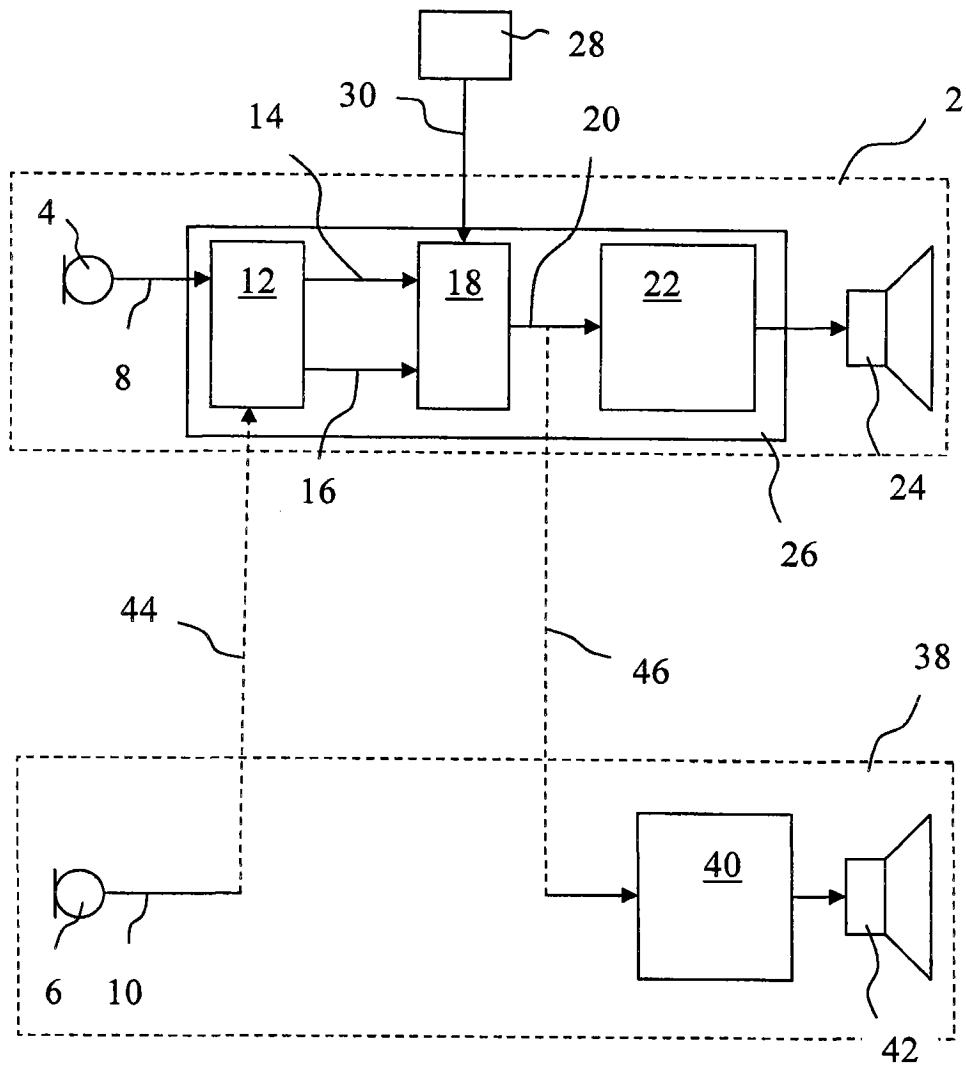


图 4

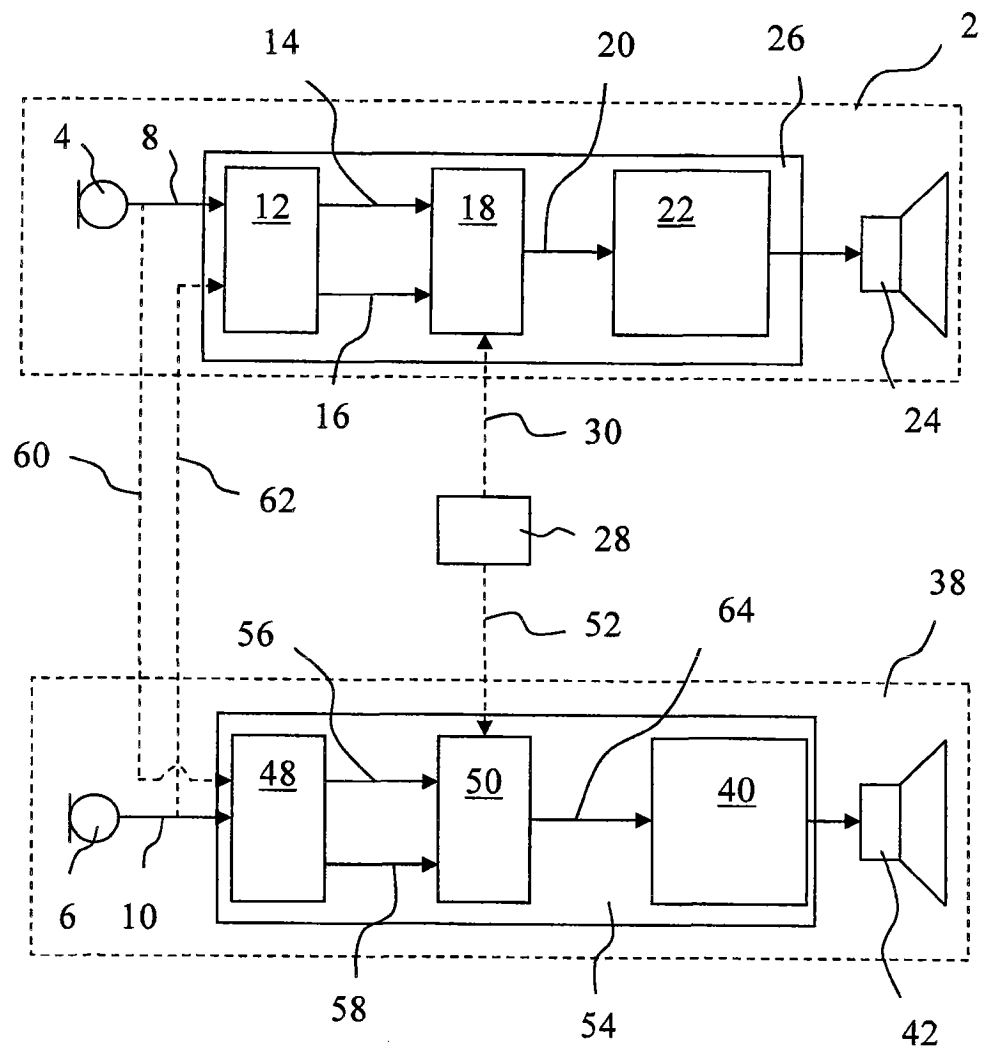


图 5

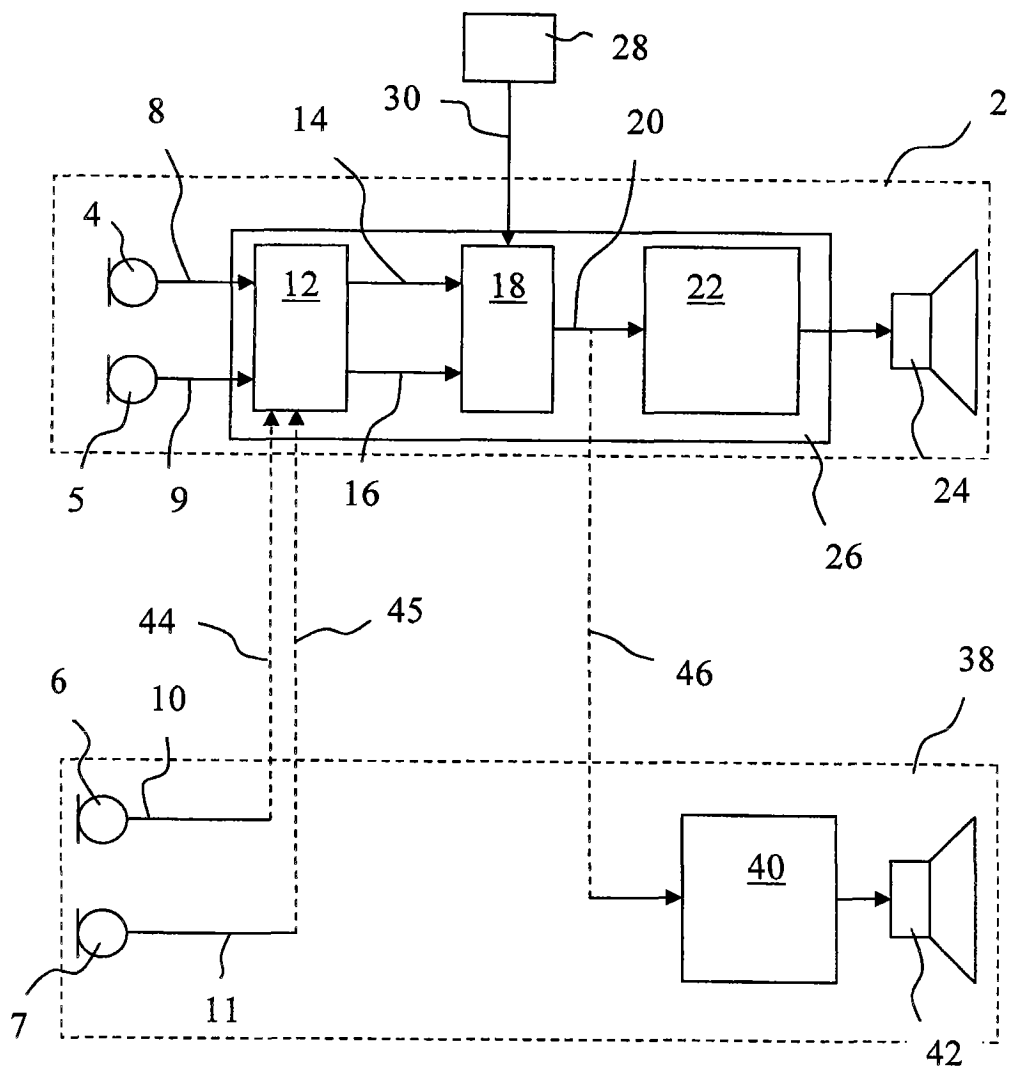


图 6

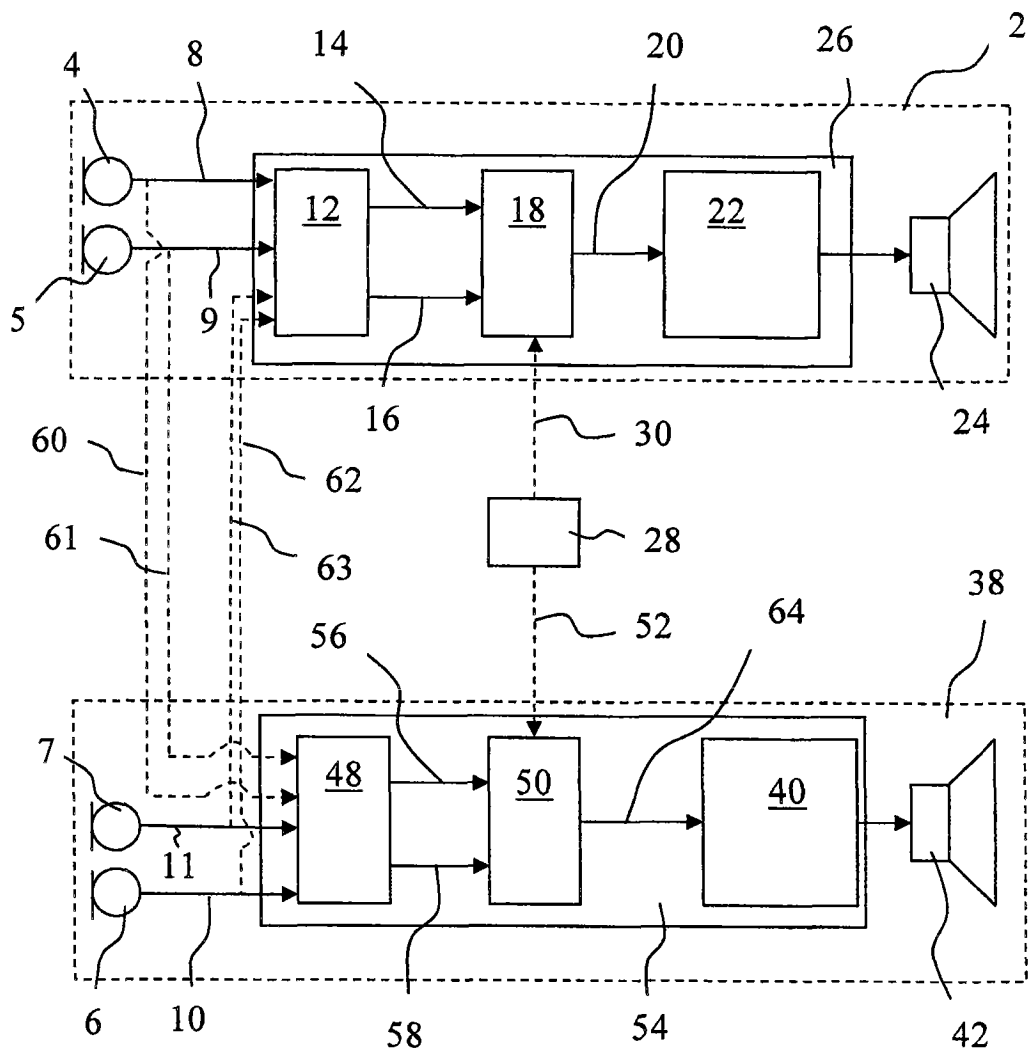


图 7

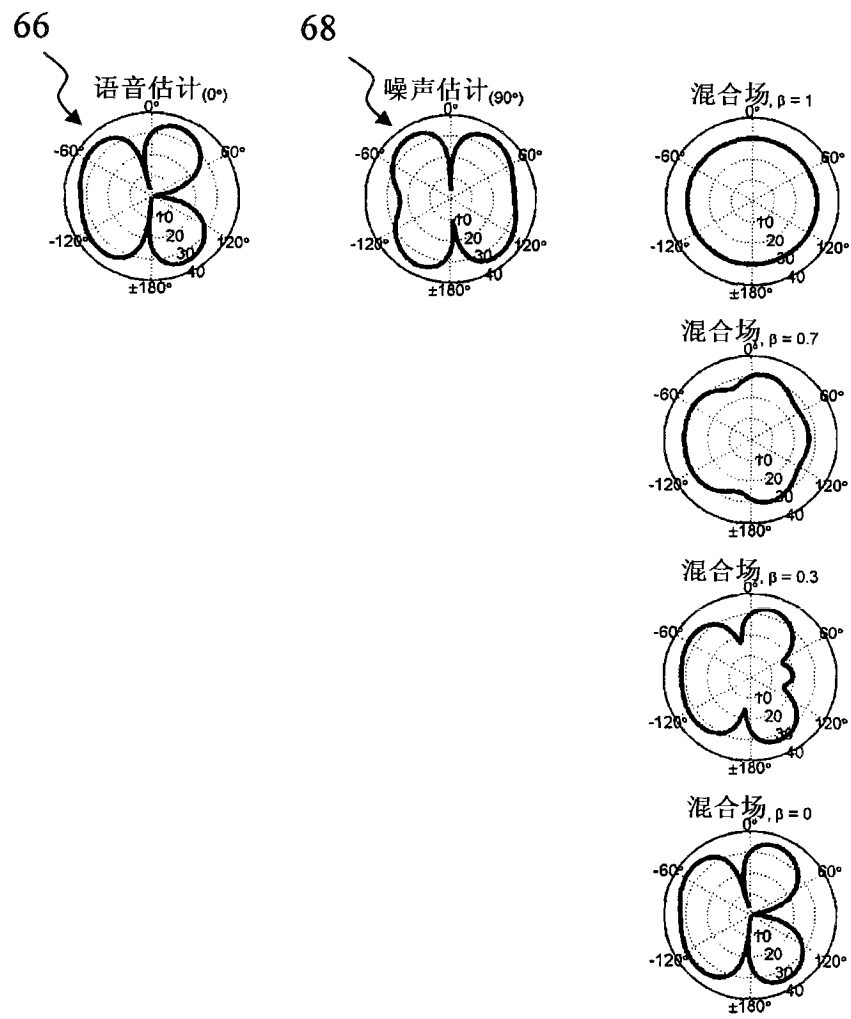


图 8A

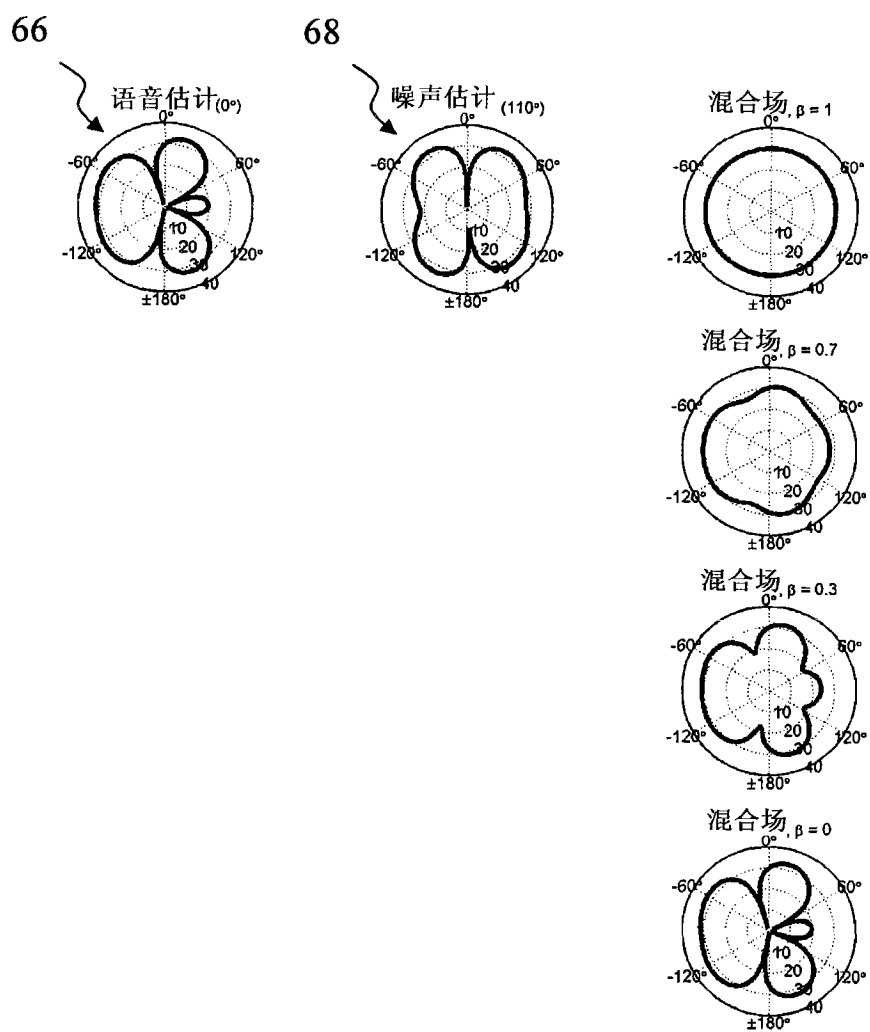


图 8B

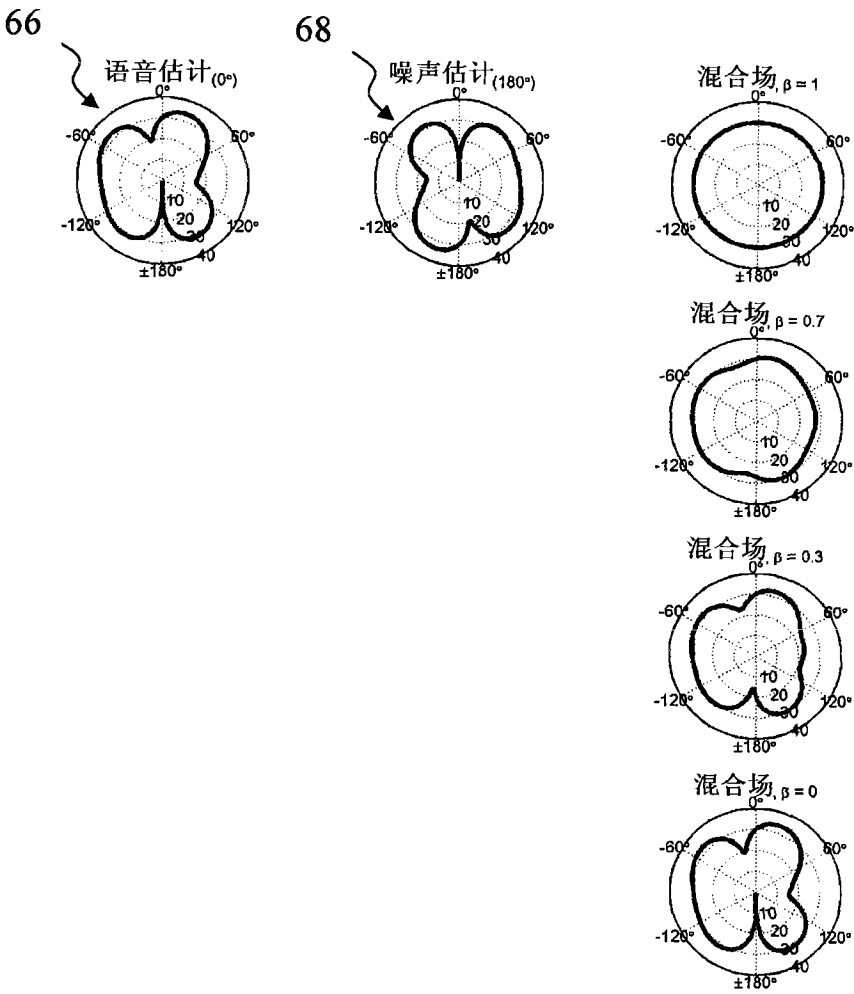


图 8C

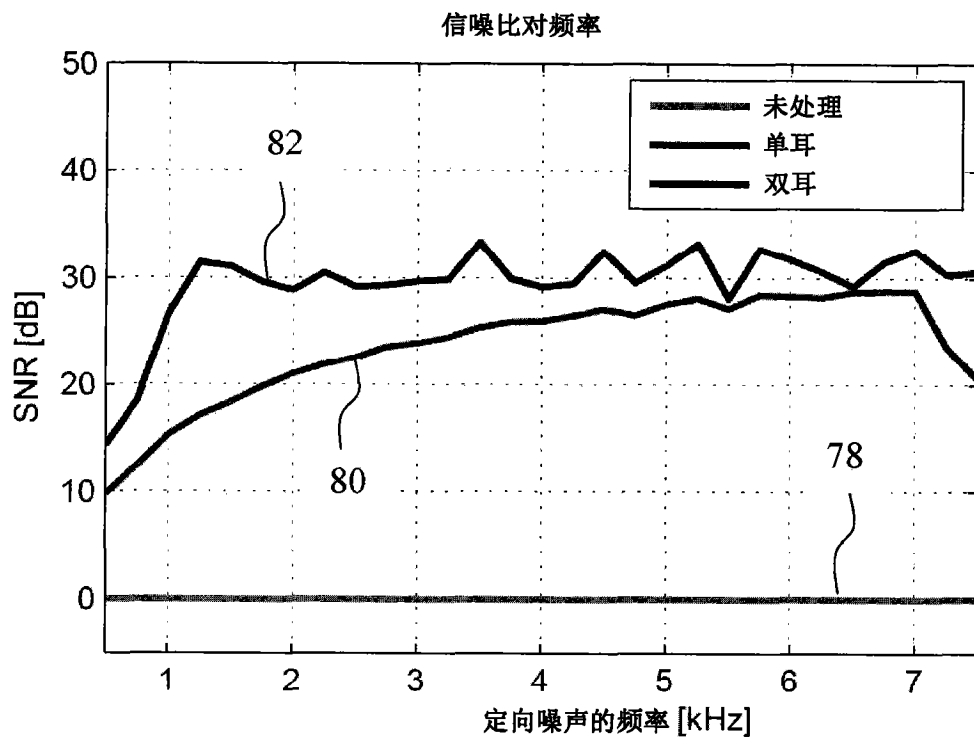


图 9

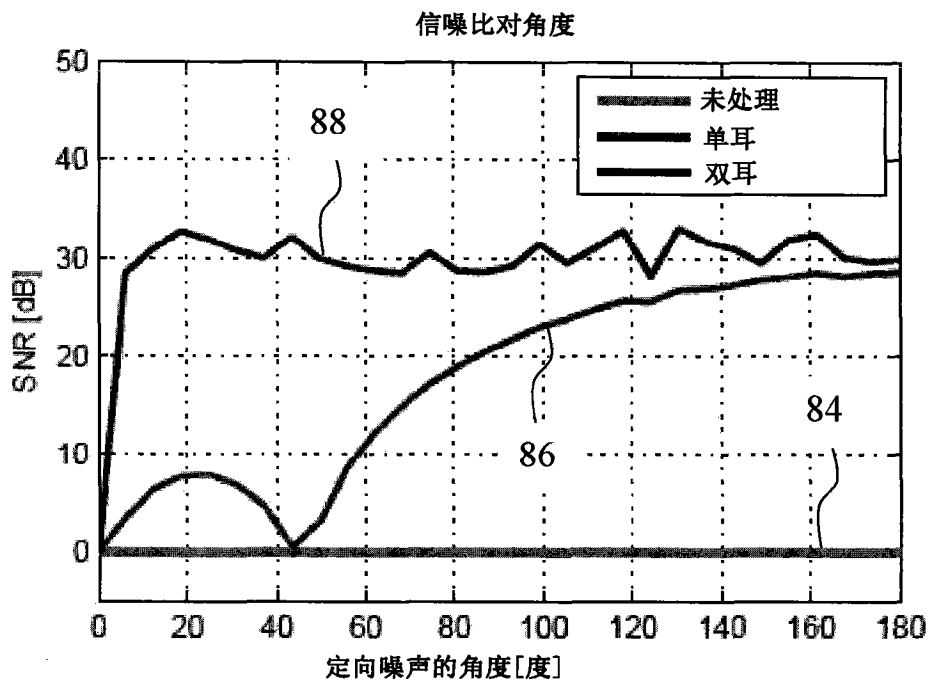


图 10

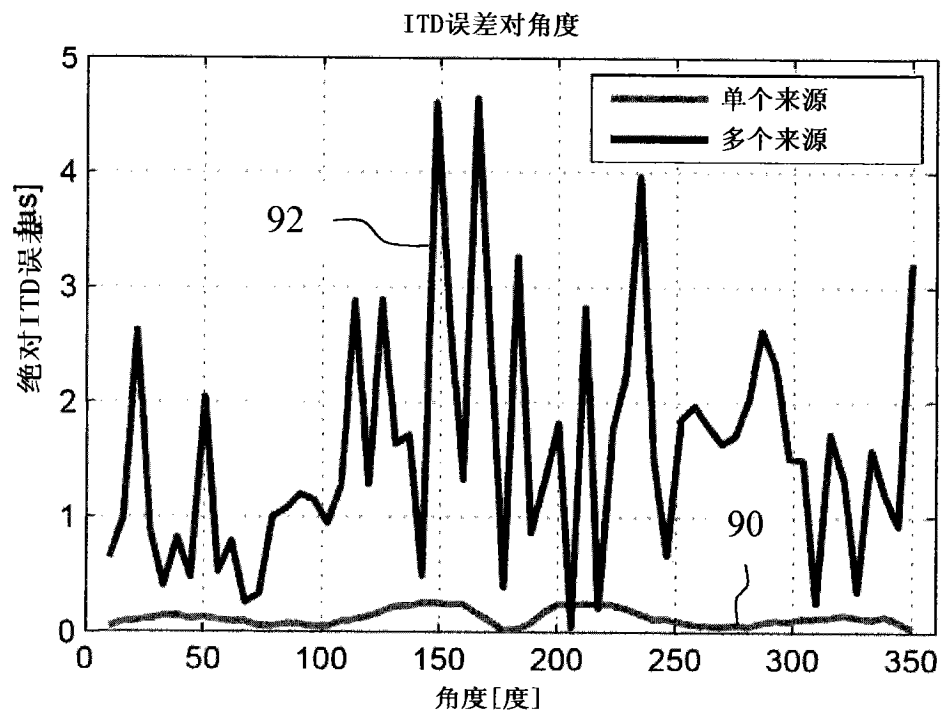


图 11

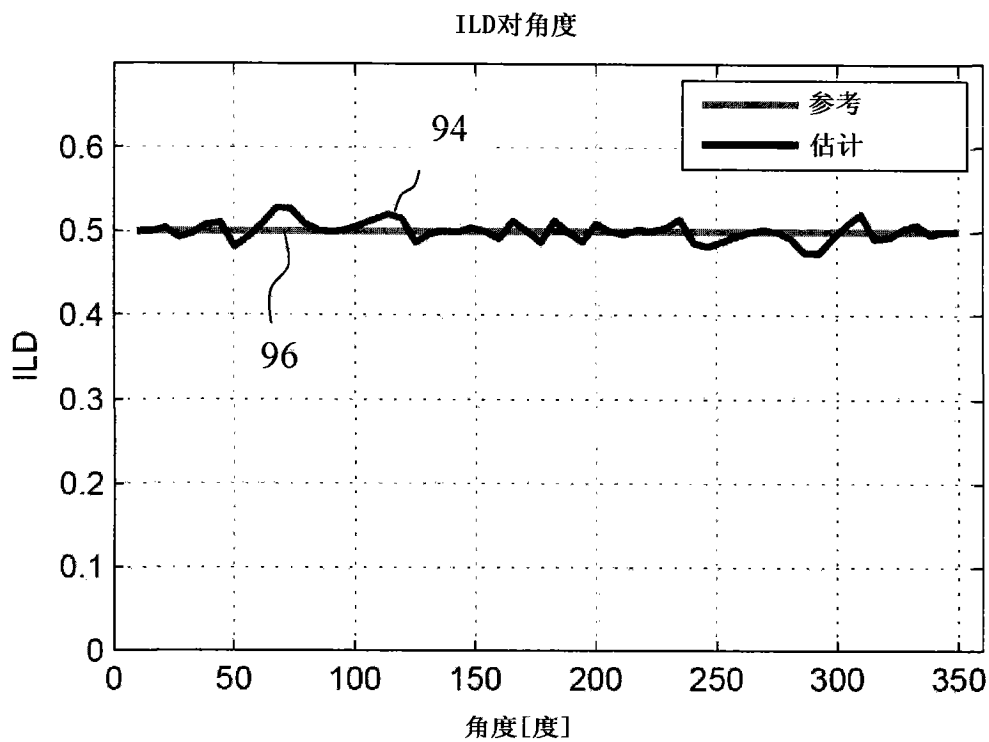


图 12