



[12] 发 明 专 利 说 明 书

[21] ZL 专利号 98812682.6

[45] 授权公告日 2004 年 3 月 24 日

[11] 授权公告号 CN 1143268C

[22] 申请日 1998.12.7 [21] 申请号 98812682.6

[30] 优先权

[32] 1997. 12. 24 [33] JP [31] 354754/1997

[86] 国际申请 PCT/JP98/05513 1998.12.7

[87] 国际公布 W099/34354 日 1999.7.8

[85] 进入国家阶段日期 2000.6.26

[71] 专利权人 三菱电机株式会社

地址 日本东京都

[72] 发明人 山浦正

审查员 杨 叁

[74] 专利代理机构 中国专利代理(香港)有限公司

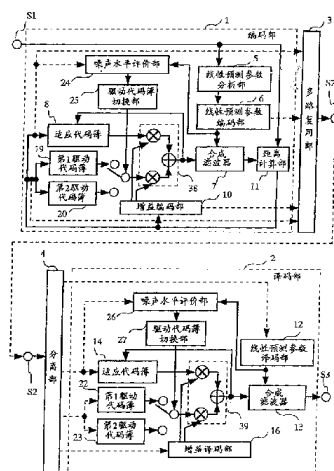
代理人 刘宗杰 叶恺东

权利要求书 1 页 说明书 11 页 附图 7 页

[54] 发明名称 声音编码方法、声音译码方法、声音编码装置和声音译码装置

[57] 摘要

在将声音信号压缩编码成数字信号的声音编码译码中，使用较少的信息量再生高品质的声音。在码驱动线性预测 (CELP) 声音编码中，使用频谱信息、功率信息和音调信息中的至少一个代码或编码结果对该编码区间内的声音的噪声水平进行评价，根据评价结果使用不同的驱动代码簿 19、20。



1. 一种声音译码方法，其特征在于：在码驱动线性预测声音译码方法中，使用频谱信息、功率信息和音调信息中的至少一个代码或译码结果，对该译码区间中的声音的噪声水平进行评价，根据评价结果使从驱动代码簿中输出的时间序列矢量的噪声水平发生变化。

2. 一种声音译码装置，其特征在于，在编码驱动线性预测声音译码装置中，包括：使用频谱信息、功率信息和音调信息中的至少一个代码或译码结果对该译码区间内的声音的噪声水平进行评价的噪声水平评价部；

根据上述噪声水平评价部的评价结果使从驱动代码簿中输出的时间序列矢量的噪声水平发生变化的噪声度控制部。

3. 一种声音译码方法，其特征在于：在码驱动线性预测声音译码方法中，使用功率信息代码或译码结果，对该译码区间中的声音的噪声水平进行评价，根据评价结果使从驱动代码簿中输出的时间序列矢量的噪声水平发生变化。

4. 一种声音译码装置，其特征在于，在编码驱动线性预测声音译码装置中，包括：使用功率信息代码或译码结果对该译码区间内的声音的噪声水平进行评价的噪声水平评价部；

根据上述噪声水平评价部的评价结果使从驱动代码簿中输出的时间序列矢量的噪声水平发生变化的噪声度控制部。

声音编码方法、声音译码方法、
声音编码装置和声音译码装置

5 技术领域

本发明涉及对声音信号进行数字信号的压缩编码译码时使用的声音编码译码方法和声音编码译码装置，特别涉及用来使用低比特率再生高品质的声音的声音编码方法、声音译码方法、声音编码装置和声音译码装置。

10 背景技术

过去，作为高效率声音编码方法，典型的有码驱动线性预测编码 (Code-Excited Linear Prediction: CELP)，对该技术，“Code-Excited Linear Prediction (CELP): High-quality speech at very low bit rates” (M. R. Shroeder and B. S. Atal 著、ICASSP'85, pp. 937-940, 15 1985) 已有叙述。

图 6 是表示一例 CELP 声音编码方法的整体构成的图。图中 101 是编码部，102 是译码部，103 是多路复用装置，104 是分离装置。编码部 101 由线性预测参数分析装置 105、线性预测参数编码装置 106、合成滤波器 107、适应代码簿 108、驱动代码簿 109、增益编码装置 110、距离计算装置 111 和加权相加计算装置 138 构成。此外，译码部 102 20 由线性预测参数译码装置 112、合成滤波器 113、适应代码簿 114、驱动代码簿 115、增益译码装置 116 和加权相加计算装置 139 构成。

在 CELP 声音编码中，将 5~50ms 作为一帧，将该帧的声音分成频谱信息和声音源信息后进行编码。首先，说明 CELP 声音编码方法的动作。在编码部 101 中，线性预测参数分析装置 105 分析输入声音 S101，25 抽出作为声音频谱信息的线性预测参数。线性预测参数编码装置 106 对该线性预测参数进行编码，将该编码后的线性预测参数作为合成滤波器的系数来设定。

其次，说明声音源信息的编码。在适应代码簿 108 中，存储过去的驱动声音源信号，并与距离计算装置 111 输入的适应代码对应输出 30 周期性的重复过去的驱动声音源信号的时间序列矢量。在驱动代码簿 109 中，存储多个时间序列矢量，该时间序列矢量构成为例如能够进行

学习,使学习用声音和它的编码声音的失真很小。从适应代码簿 108、驱动代码簿 109 来的各时间序列矢量与增益编码装置 110 给出的各增益对应,在加权相加计算装置 138 中进行加权相加,将该计算结果作为驱动声音信号供给合成滤波器 107,得到编码声音。距离计算装置 111 5 求出编码声音和输入声音 S101 的距离,寻求距离最小的适应代码、驱动代码和增益。在上述编码结束后,将线性预测参数的代码以及使输入声音和编码声音的失真最小的适应代码、驱动代码、增益的代码作为编码结果输出。

其次,说明 CPEL 声音译码方法的动作。

10 另一方面,在声音译码部 102 中,线性预测参译编码装置 112 根据线性预测参数的代码对该线性预测参数进行译码,并作为合成滤波器的系数来设定。其次,适应代码簿 114 与适应代码对应输出周期性的重复过去的驱动声音源信号的时间序列矢量,驱动代码簿 115 与驱动代码对应时间序列矢量。这些时间序列矢量与增益译码装置 15 中从增益代码译码的各增益对应,在加权相加计算装置 139 中进行加权相加,将该计算结果作为驱动声音信号供给合成滤波器 113,得到输出声音 S103。

此外,在 CELP 声音编码译码方法中,作为以提高再生声音品质为目的进行改良的先有的声音编码译码方法,有“Phonetically-based vector excitation coding of speech at 3.6kbps”(S.wang and A.Gersho 著、ICASSP'89, pp.49-52,1989)所示的方法。图 7 20 示出一例该先有的声音编码译码方法的整体构成,对与图 6 对应的装置添加相同的符号,在图中的编码部 101 中,117 是声音状态判定装置,118 是驱动代码簿切换装置,119 是第 1 驱动代码簿,120 是第 2 驱动代码簿。此外,在图中的译码装置 102 中,121 是驱动代码簿切换装置,122 是第 1 驱动代码簿,123 是第 2 驱动代码簿。说明这样构成的编码译码方法的动作。首先,在编码装置 101 中,声音状态判定装置 117 分析输入声音 S101,判定声音状态例如有声、25 无声两种状态中的哪一种状态。驱动代码簿切换装置 118 根据该声音状态的判定结果切换驱动代码簿,例如,若是有声则使用第 1 驱动代码簿 119 编码,若是无声则使用第 2 驱动代码簿 120 编码,此外,对使用了哪一个驱动代码簿也进行编码。30

其次，在译码装置 102 中，驱动代码簿切换装置 121 与在编码装置中使用了哪一个驱动代码簿的代码对应切换到第 1 驱动代码簿或第 2 驱动代码簿，使其与编码装置 101 使用的驱动代码簿相同。通过这样的构成，对声音的每一个状态准备一个与编码适应的驱动代码簿，通过
5 与输入的声音状态对应切换使用驱动代码簿，可以提高再生声音的品质。

此外，作为不增加比特数去切换多个驱动代码簿的先有的声音编码译码方法，有特开平 8—185198 号公报公开的方法。它是与用适应代码簿选择的音调周期对应去切换使用多个驱动代码簿的方法。因此，
10 可以在不增加传送信息的情况下使用与输入信号的特征相适应的驱动代码簿。

如上所述，在图 6 所示的先有的声音编码译码方法中，使用单一的驱动代码簿生成合成声音。为了即使在低比特率时也能得到高品质的编码声音，存储在驱动代码簿中的时间序列矢量变成包含很多脉冲的无噪声的东西。因此，当将背景噪声或磨擦性子音等有噪声的声音
15 编码合成时，编码声音存在产生“叽哩叽哩”“噉哩噉哩”等不自然的聲音的问题。若使驱动编码簿只由带噪声的时间序列矢量构成，虽然可以解决该问题，但作为编码声音的整体品质却变差了。

此外，在已改良的图 7 所示的先有的声音编码译码方法中，与输入声音的状态对应切换多个驱动代码簿并生成编码声音。因此，对例如输入声音是有噪声的无声部分，可以使用由有噪声的时间序列矢量构成的驱动代码簿，对除此之外的有声部分可以使用由无噪声的时间序列矢量构成的驱动代码簿，即使对有噪声的声音进行编码、也不会
20 发生“叽哩叽哩”的声音。但是，因译码侧也使用和编码侧相同的驱动代码簿，故有必要对使用了哪一个驱动编码簿的信息重新进行编码传送，存在妨碍低比特率化的问题。

此外，在不增加发送比特数的情况下切换多个驱动代码簿的先有的声音编码译码方法中，与用适应代码选择的音调周期对应切换驱动代码簿。但是，因用适应代码选择的音调周期与实际的声音音调周期
30 有差别，只根据该值不能判定输入声音的状态是有噪声还是无噪声，故不能解决声音的噪声部分的编码声音不自然的问题。

发明内容

本发明是为了解决有关的问题而提出的，其目的在于提供一种声音编码译码方法和声音编码译码装置，即使在低比特率的情况下也能再生高品质的声音。

- 5 本发明的声音译码方法，其特征在于：在码驱动线性预测声音译码方法中，使用频谱信息、功率信息和音调信息中的至少一个代码或译码结果，对该译码区间中的声音的噪声水平进行评价，根据评价结果使从驱动代码簿中输出的时间序列矢量的噪声水平发生变化。

- 10 本发明的声音译码装置，其特征在于，在编码驱动线性预测声音译码装置中，包括：使用频谱信息、功率信息和音调信息中的至少一个代码或译码结果对该译码区间内的声音的噪声水平进行评价的噪声水平评价部；

根据上述噪声水平评价部的评价结果使从驱动代码簿中输出的时间序列矢量的噪声水平发生变化的噪声度控制部。

- 15 本发明的又一种声音译码方法，其特征在于：在码驱动线性预测声音译码方法中，使用功率信息代码或译码结果，对该译码区间中的声音的噪声水平进行评价，根据评价结果使从驱动代码簿中输出的时间序列矢量的噪声水平发生变化。

- 20 本发明的又一种声音译码装置，其特征在于，在编码驱动线性预测声音译码装置中，包括：使用功率信息代码或译码结果对该译码区间内的声音的噪声水平进行评价的噪声水平评价部；

根据上述噪声水平评价部的评价结果使从驱动代码簿中输出的时间序列矢量的噪声水平发生变化的噪声度控制部。

附图的简单说明

- 25 图 1 是表示本发明的声音编码和声音译码装置的实施形态 1 的整体构成的方框图。

图 2 是向图 1 的实施形态 1 的噪声水平评价的说明提供的表。

图 3 是表示本发明的声音编码和声音译码装置的实施形态 3 的整体构成的方框图。

- 30 图 4 是表示本发明的声音编码和声音译码装置的实施形态 5 的整体构成的方框图。

图 5 是向图 4 的实施形态 5 的加权决定处理的说明提供的表。

图 6 是表示先有的 CELP 声音编码译码装置的整体构成的方框图。

图 7 是表示过去改良了的 CELP 声音编码译码装置的整体构成的方框图。

发明的具体实施方式下面，参照附图说明本发明的实施形态。

实施形态 1.

5 图 1 示出本发明的声音编码方法和声音译码方法的实施形态 1 的整体构成的方框图。图中，1 是编码部，2 是译码部，3 是多路复用部，4 是分离部。编码部 1 由线性预测参数分析部 5、线性预测参数编码部 6、合成滤波器 7、适应代码簿 8、增益编码部 10、距离计算装置 11、第 1 驱动代码簿 19、第 2 驱动代码簿 20、噪声水平评价部 24、驱动代码簿切换部 25 和加权相加计算部 38 构成。此外，译码部 2 由线性预测参数译码部 12、合成滤波器 13、适应代码簿 14、第 1 驱动代码簿 22、第 2 驱动代码簿 23、噪声水平评价部 26、驱动代码簿切换部 27、增益译码部 16 和加权相加计算部 39 构成。图 1 中的 5 是作为频谱信息分析部的线性预测参数分析部，分析输入声音 S1，抽出作为声音频谱信息的线性预测参数，6 是作为频谱信息编码部的线性预测参数编码部，
10 对作为频谱信息的该线性预测参数进行编码，将该编码后的线性预测参数作为合成滤波器 7 的系数来设定，19、22 是存储多个无噪声的时间序列矢量的第 1 驱动代码簿，20、23 是存储多个有噪声的时间序列矢量的第 2 驱动代码簿，24、26 是评价噪声水平的噪声水平评价部，25、
15 27 是根据噪声水平切换驱动代码簿的驱动代码簿切换部。

下面，说明动作。首先，在编码部 1 中，线性预测参数分析部 5 分析输入声音 S1，抽出作为声音频谱信息的线性预测参数。线性预测参数编码部 6 对该线性预测参数进行编码，将该编码后的线性预测参数作为合成滤波器 7 的系数来设定，同时，向噪声水平评价部 24 输出。
20 其次，说明声音源信息的编码。适应代码簿 8 存储过去的驱动声音源信号，并与距离计算装置 11 输入的适应代码对应输出周期性的重复过去的驱动声音源信号的时间序列矢量。噪声水平评价部 24 根据从上述线性预测参数编码部 6 输入的已编码的线性预测参数和适应代码，例如如图 2 所示那样，从频谱的倾斜、短期预测增益和音调变动去评价
25 该编码区间的噪声水平，并将评价结果输出给驱动代码簿切换部 25。驱动代码簿切换部 25 根据上述噪声水平的评价结果去切换编码时用的驱动代码簿，例如，若噪声水平低，则切换到第 1 驱动代码簿 19，若
30

噪声水平高，则切换到第 2 驱动代码簿 20。

在第 1 驱动代码簿 19 中存储多个无噪声的时间序列矢量，该时间序列矢量构成为例如能够进行学习，使学习用声音和它的编码声音的失真很小。此外，在第 2 驱动代码簿 20 中存储多个有噪声的时间序列矢量，例如，存储由随机噪声生成的多个时间序列矢量，输出与从距离计算部 11 输入的各个驱动代码对应的的时间序列矢量。从适应代码簿 8、第 1 驱动代码簿 19 或第 2 驱动代码簿 20 来的各时间序列矢量与增益编码部 10 加给的各增益对应，在加权相加计算部 38 中进行加权相加，将该计算结果作为驱动声音信号供给合成滤波器 7，得到编码声音。距离计算部 11 求出编码声音和输入声音 S1 的距离，寻求距离最小的适应代码、驱动代码和增益。在上述编码结束后，将线性预测参数的代码以及使输入声音和编码声音的失真最小的适应代码、驱动代码、增益的代码作为编码结果输出。以上是本实施形态 1 的声音编码方法的特征动作。

其次，说明译码部 2。在译码部 2 中，线性预测参数译码部 12 从线性预测参数的代码中译码出线性预测参数并作为合成滤波器 13 的系数来设定，同时，向噪声水平评价部 26 输出。其次，说明声音源信息的译码。适应代码簿 14 与适应代码对应，输出周期地重复过去的驱动声音源信号的时间序列矢量。噪声水平评价部 26 使用和编码部 1 的噪声水平评价部 24 相同的方法，根据从上述线性预测参数译码部 12 输入的已译码的线性预测参数和适应代码去评价噪声水平，并将评价结果输出给驱动代码簿切换部 27。驱动代码簿切换部 27 和编码部 1 的驱动代码簿切换部 25 一样，根据上述噪声水平的评价结果切换第 1 驱动代码簿 22 和第 2 驱动代码簿 23。

在第 1 驱动代码簿 22 中存储多个无噪声的时间序列矢量，该时间序列矢量构成为例如能够进行学习，使学习用声音和它的编码声音的失真很小，而在第 2 驱动代码簿 20 中存储多个有噪声的时间序列矢量，例如，存储由随机噪声生成的多个时间序列矢量，输出与从距离计算部 11 输入的各个驱动代码对应的的时间序列矢量。从适应代码簿 14 和第 1 驱动代码簿 22 或第 2 驱动代码簿 23 来的各时间序列矢量与在增益译码部 16 中从增益代码译码出的各增益对应，在加权相加计算部 39 中进行加权相加，将该计算结果作为驱动声音信号

供给合成滤波器 13, 得到输出声音 S3。以上是本实施形态 1 的声音译码方法的特征动作。

若按照该实施形态 1, 通过根据代码和编码结果对输入声音的噪声水平进行评价并根据评价结果使用不同的驱动代码簿, 可以用少量的信息再生出高品质的声音。

此外, 在上述实施形态中, 对驱动代码簿 19、20、22、23 说明了存储多个时间序列矢量的情况, 但只要存储至少一个时间序列矢量, 就可以实施本发明。

实施形态 2

在上述实施形态 1 中, 切换使用两个驱动代码簿, 但也可以具有三个以上的驱动代码簿, 根据噪声水平进行切换使用。若按照该实施形态 2, 因为不只是将声音分成有噪声和无噪声两种类型, 对于有一点噪声的中间状态的声音也可以使用与其相应的驱动代码簿, 所以能够再生出高品质的声音。

实施形态 3

图 3 示出本发明的声音编码方法和声音译码方法的实施形态 3 的整体构成, 对与图 1 对应的部分添加相同的符号, 图中 28、30 是存储有噪声的时间序列矢量的驱动代码簿, 29、31 是将时间序列矢量的小振幅样品的振幅值为零的样品抽值部。

下面, 说明动作。首先, 在编码部 1 中, 线性预测参数分析部 5 分析输入声音 S1, 抽出作为声音频谱信息的线性预测参数。线性预测参数编码部 6 对该线性预测参数进行编码, 将该编码后的线性预测参数作为合成滤波器 7 的系数来设定, 同时, 向噪声水平评价部 24 输出。其次, 说明声音源信息的编码。适应代码簿 8 存储过去的驱动声音源信号, 并与距离计算部 11 输入的适应代码对应输出周期性的重复过去的驱动声音源信号的时间序列矢量。噪声水平评价部 24 根据从上述线性预测参数编码部 6 输入的已编码的线性预测参数和适应代码, 例如从频谱的倾斜、短期预测增益和音调变动去评价该编码区间的噪声水平, 并将评价结果输出给样品抽值部 29。

在驱动代码簿 28 中存储例如由随机噪声生成的多个时间序列矢量, 输出与从距离计算部 11 输入驱动代码对应的时间序列矢量。样品抽值部 29 根据上述噪声水平的评价结果, 若噪声水平低, 则在从上述

驱动代码簿 28 输入的时间序列矢量中输出使例如未达到规定的振幅值的样品的振幅值为零的时间序列矢量，此外，若噪声水平高，则直接输出从上述驱动代码簿 28 输入的时间序列矢量。从适应代码簿 8、样品抽值部 29 来的各时间序列矢量与增益编码部 10 加给的各增益对应，在加权相加计算部 38 中进行加权相加，将该计算结果作为驱动声音信号供给合成滤波器 7，得到编码声音。距离计算部 11 求出编码声音和输入声音 S1 的距离，寻求距离最小的适应代码、驱动代码和增益。在上述编码结束后，将线性预测参数的代码以及使输入声音和编码声音的失真最小的适应代码、驱动代码、增益的代码作为编码结果 S2 输出。以上是本实施形态 1 的声音编码方法的特征动作。

其次，说明译码部 2。在译码部 2 中，线性预测参数译码部 12 从线性预测参数的代码中译码出线性预测参数并作为合成滤波器 13 的系数来设定，同时，向噪声水平评价部 26 输出。其次，说明声音源信息的译码。适应代码簿 14 与适应代码对应，输出周期地重复过去的驱动声音源信号的时间序列矢量。噪声水平评价部 26 使用和编码部 1 的噪声水平评价部 24 相同的方法，根据从上述线性预测参数译码部 12 输入的已译码的线性预测参数和适应代码去评价噪声水平，并将评价结果输出给样品抽值部 31。

驱动代码簿 30 与驱动代码对应输出时间序列矢量。样品抽值部 31 通过和上述编码部 1 的样品抽值部 29 同样的处理，根据上述噪声评价结果输出时间序列矢量。从适应代码簿 14 和样品抽值部 31 来的各时间序列矢量与增益译码部 16 加给的各增益对应，在加权相加计算部 39 中进行加权相加，将该计算结果作为驱动声音源信号供给合成滤波器 13，得到输出声音 S3。

若按照该实施形态 3，具有存储有噪声的时间序列矢量的驱动代码簿，通过根据声音的噪声水平的结果对驱动声音源的信息样品进行抽值来生成噪声水平低的驱动声音源，可以用少量的信息再生出高品质的声音。此外，因不需要多个驱动代码簿，故具有能够减少用于存储驱动代码簿的存储器的数量的效果。

实施形态 4

在上述实施形态 3 中，对时间序列矢量的样品有抽值和不抽值两种选择，但也可以在抽值样品时根据噪声水平变更振幅阈值。若按照

该实施形态 4，因为不只是将声音分成有噪声和无噪声两种类型，对于有一点噪声的中间状态的声音也可以生成并使用与其相应的时间序列矢量，所以能够再生出高品质的声音。

实施形态 5

5 图 4 示出本发明的声音编码方法和声音译码方法的实施形态 5 的整体构成，对与图 1 对应的部分添加相同的符号，图中 32、35 是存储有噪声的时间序列矢量的第 1 驱动代码簿，33、36 是存储无噪声的时间序列矢量的第 2 驱动代码簿，34、37 是权重决定部。

下面，说明动作。首先，在编码部 1 中，线性预测参数分析部 5
10 分析输入声音 S1，抽出作为声音频谱信息的线性预测参数。线性预测参数编码部 6 对该线性预测参数进行编码，将该编码后的线性预测参数作为合成滤波器 7 的系数来设定，同时，向噪声水平评价部 24 输出。其次，说明声音源信息的编码。适应代码簿 8 存储过去的驱动声音源信号，并与距离计算部 11 输入的适应代码对应输出周期性的重复过去的
15 的驱动声音源信号的时间序列矢量。噪声水平评价部 24 根据从上述线性预测参数编码部 6 输入的已编码的线性预测参数和适应代码，例如从频谱的倾斜、短期预测增益和音调变动去评价该编码区间的噪声水平，并将评价结果输出给权重决定部 34。

在第 1 驱动代码簿 32 中存储例如由随机噪声生成的多个有噪声的时间序列矢量，输出与驱动代码对应的时间序列矢量。在第 2 驱动代码簿 20 中存储多个时间序列矢量，该时间序列矢量构成为例如能够进行学习，使学习用声音和它的编码声音的失真很小。输出与从距离计算部 11 输入的驱动代码对应的时间序列矢量。重量决定部 34 根据从上述噪声水平评价部 24 输入的噪声水平评价结果，例如按照图 5 决定
25 加给第 1 驱动代码簿 32 的时间序列矢量和第 1 驱动代码簿 32 的时间序列矢量的权重。第 1 驱动代码簿 32 和第 2 驱动代码簿 33 的各时间序列矢量根据上述权重决定部 34 给出的权重进行加权相加。从适应代码簿 8 输出的时间序列矢量和上述加权相加后生成的时间序列矢量与增益编码部 10 加给的各增益对应，在加权相加计算部 38 中进行加权
30 相加，将该计算结果作为驱动声音信号供给合成滤波器 7，得到编码声音。距离计算部 11 求出编码声音和输入声音 S1 的距离，寻求距离最小的适应代码、驱动代码和增益。在上述编码结束后，将线性预测参

数的代码以及使输入声音和编码声音的失真最小的适应代码、驱动代码、增益的代码作为编码结果输出。

其次，说明译码部 2。在译码部 2 中，线性预测参数译码部 12 从线性预测参数的代码中译码出线性预测参数并作为合成滤波器 13 的系数来设定，同时，向噪声水平评价部 26 输出。其次，说明声音源信息的译码。适应代码簿 14 与适应代码对应，输出周期地重复过去的驱动声音源信号的时间序列矢量。噪声水平评价部 26 使用和编码部 1 的噪声水平评价部 24 相同的方法，根据从上述线性预测参数译码部 12 输入的已译码的线性预测参数和适应代码去评价噪声水平，并将评价结果输出给权重决定部 37。

第 1 驱动代码簿 35 和第 2 驱动代码部 36 与驱动代码对应输出时间序列矢量。权重决定部 37 和编码部 1 的权重决定部 34 一样，根据从上述噪声水平评价部 26 输入的噪声水平评价结果给出权重。从第 1 驱动代码簿 35、第 2 驱动代码簿 36 来的各时间序列矢量与上述权重决定部 37 加给的各权重对应进行加权相加。从适应代码簿 14 输出的时间序列矢量和上述权重相加生成的时间序列矢量与在增益译码部 16 中从增益代码译码出的各增益对应，在加权相加计算部 39 中进行加权相加，将该计算结果作为驱动声音信号供给合成滤波器 13，得到输出声音 S3。

若按照该实施形态 5，根据代码和编码结果对输入声音的噪声水平进行评价并根据评价结果对有噪声的时间序列矢量和无噪声的时间序列矢量进行加权相加后再使用，因此，可以用少量的信息再生出高品质的声音。

实施形态 6

在上述实施形态 1~5 中，进而还可以根据噪声水平的评价结果去变更增益的代码簿。若按照该实施形态 6，因为可以根据驱动代码部使用最佳的增益代码簿，所以能够再生出高品质的声音。

实施形态 7

在上述实施形态 1~6 中，对声音的噪声水平进行评价并根据评价结果切换驱动代码簿，也可以分别对有声音的突然出现和破裂性子音等进行判定、评价并根据评价结果切换驱动代码簿。若按照该实施形态 7，因为不只对声音的噪声状态进行分类，而是对有声音的突然出现

和破裂性子音等进一步进行仔细分类，可以使用各自合适的驱动代码部，所以能够再生出高品质的声音。

实施形态 8

在上述实施形态 1~6 中，从图 2 所示的频谱倾斜、短期预测增益和音调变动去评价编码区间的噪声水平，但也可以使用相对适应代码簿的输出的增益值的大小去进行评价。

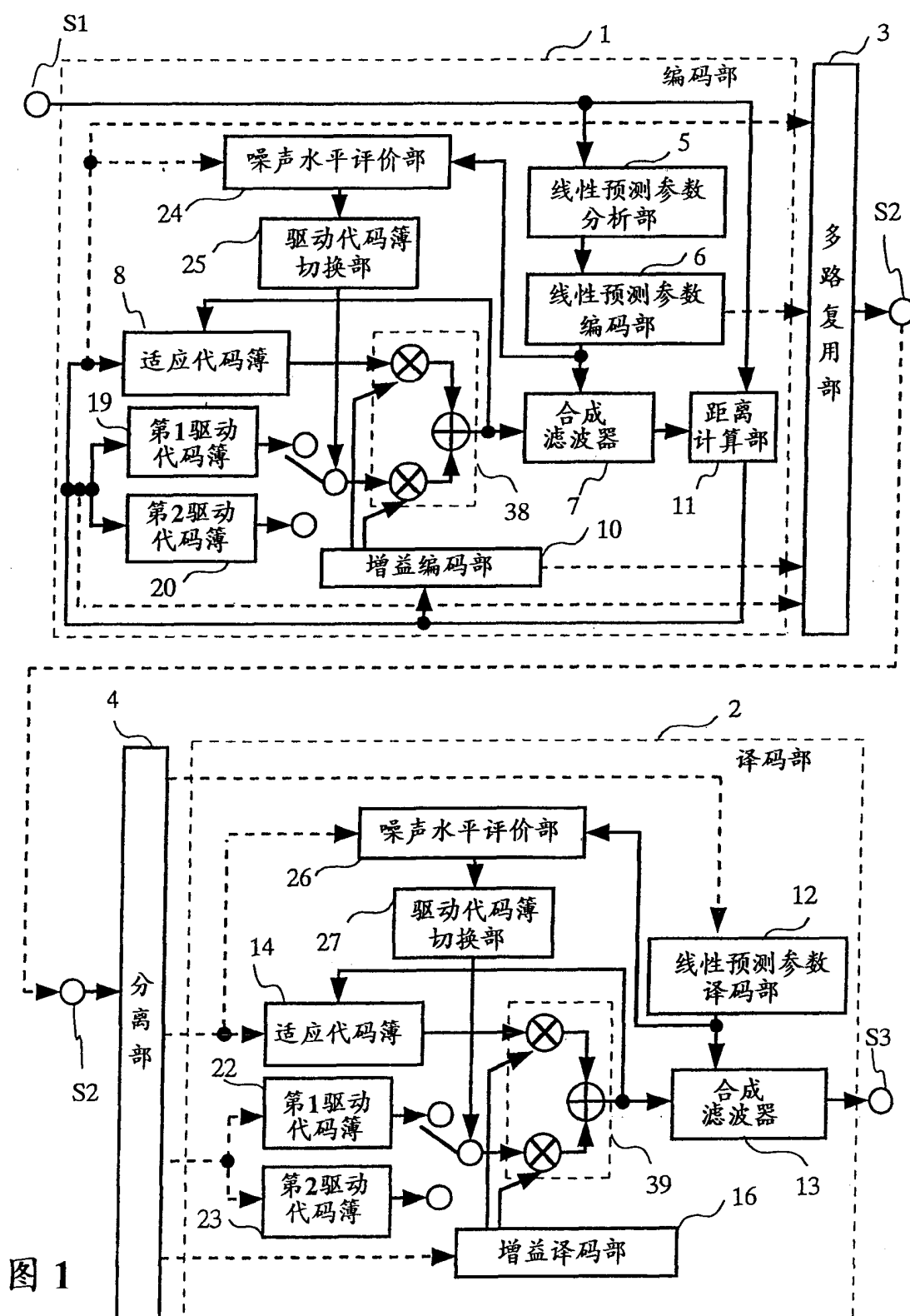
若按照本发明的声音编码方法和声音译码方法以及声音编码装置和声音译码装置，使用频谱信息、功率信息和音调信息中的至少一个代码或编码结果去评价该编码区间的噪声水平，并根据评价结果使用不同的驱动代码簿，所以，能用少量的信息再生高品质的声音。

此外，若按照本发明的声音编码方法和声音译码方法，具有多个驱动代码簿，所存储的驱动声音源的噪声水平不同，根据声音的噪声水平的评价结果，切换使用多个驱动代码簿，所以，能用少量的信息再生高品质的声音。

此外，若按照本发明的声音编码方法和声音译码方法，根据声音的噪声水平的评价结果，使存储在驱动代码簿中的时间序列矢量的噪声水平变化，所以，能用少量的信息再生高品质的声音。

此外，若按照本发明的声音编码方法和声音译码方法，具有存储有噪声的时间序列矢量的驱动代码簿，根据声音的噪声水平的评价结果，通过抽值时间序列矢量的信息样品去生成噪声水平低的时间序列矢量，所以，能用少量的信息再生高品质的声音。

此外，若按照本发明的声音编码方法和声音译码方法，具有存储有噪声的时间序列矢量的第 1 驱动代码簿和存储无噪声的时间序列矢量的第 2 驱动代码簿，根据声音的噪声水平的评价结果，对第 1 驱动代码簿的时间序列矢量和第 2 驱动代码簿的时间序列矢量进行加权相加并生成时间序列矢量，所以，能用少量的信息再生高品质的声音。



噪声水平	小 ←————→ 大
频谱倾斜	低频区倾斜 ←————→ 平坦、高频区倾斜
短期预测增益	大 ←————→ 小
音调变动	小 ←————→ 大

图 2

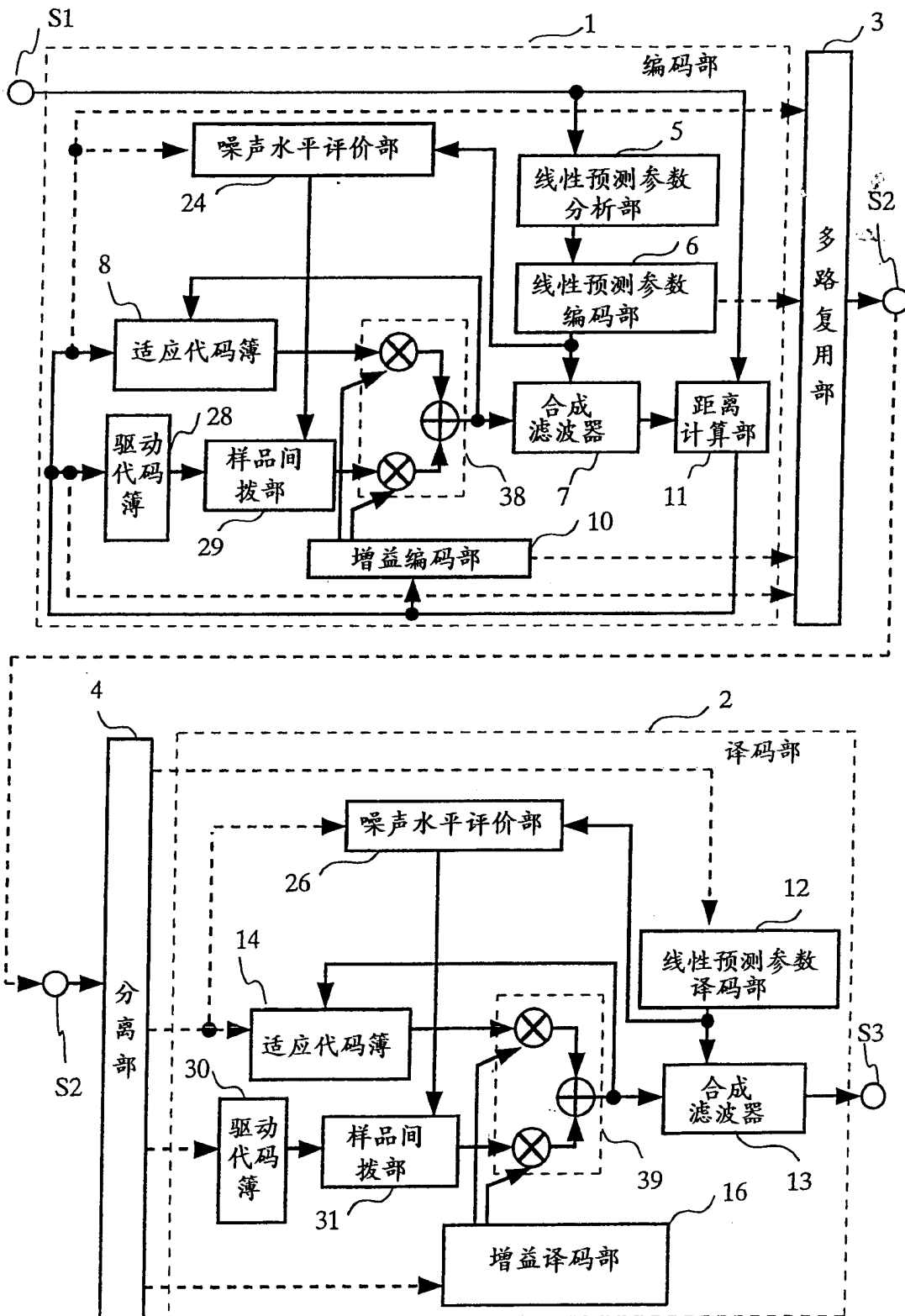


图 3

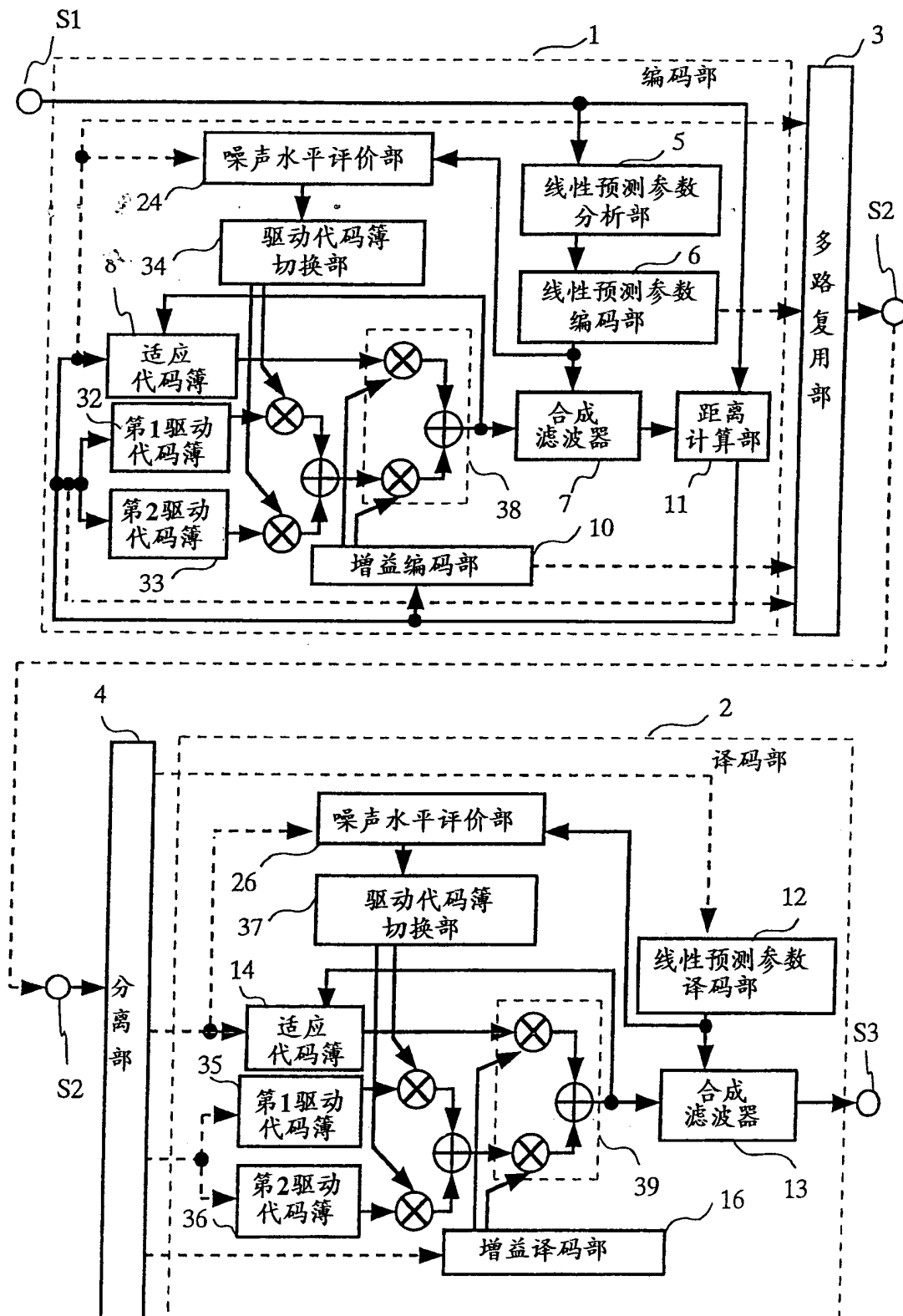


图 4

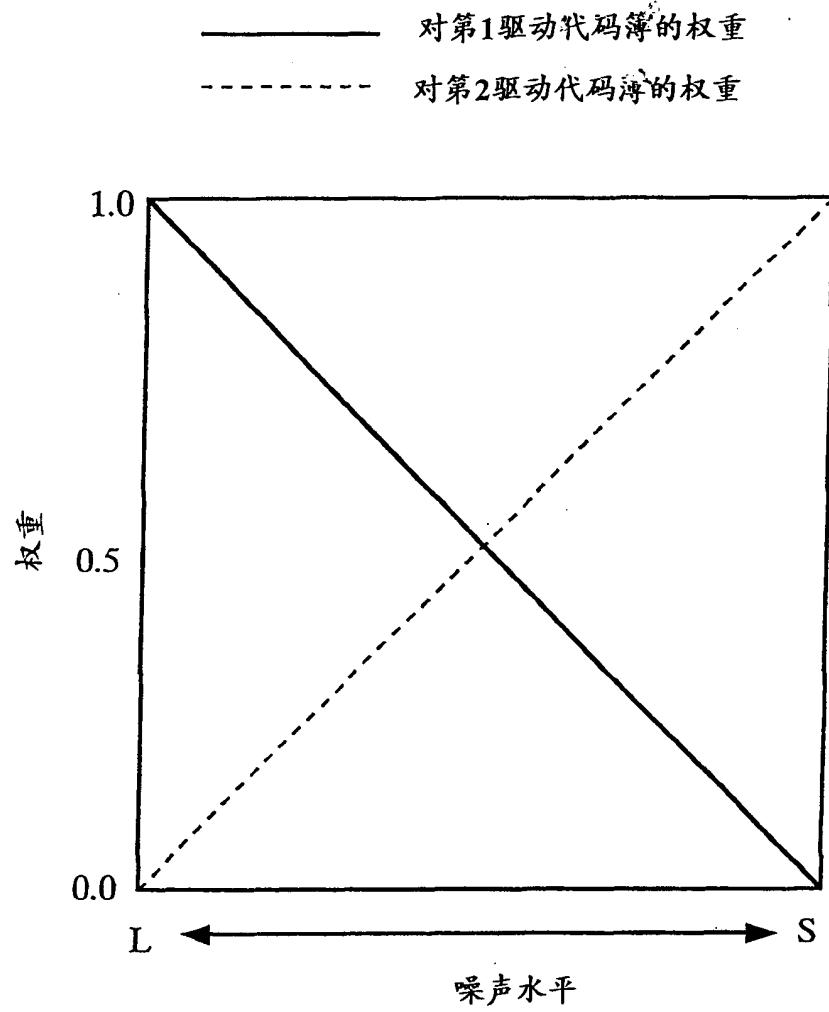


图 5

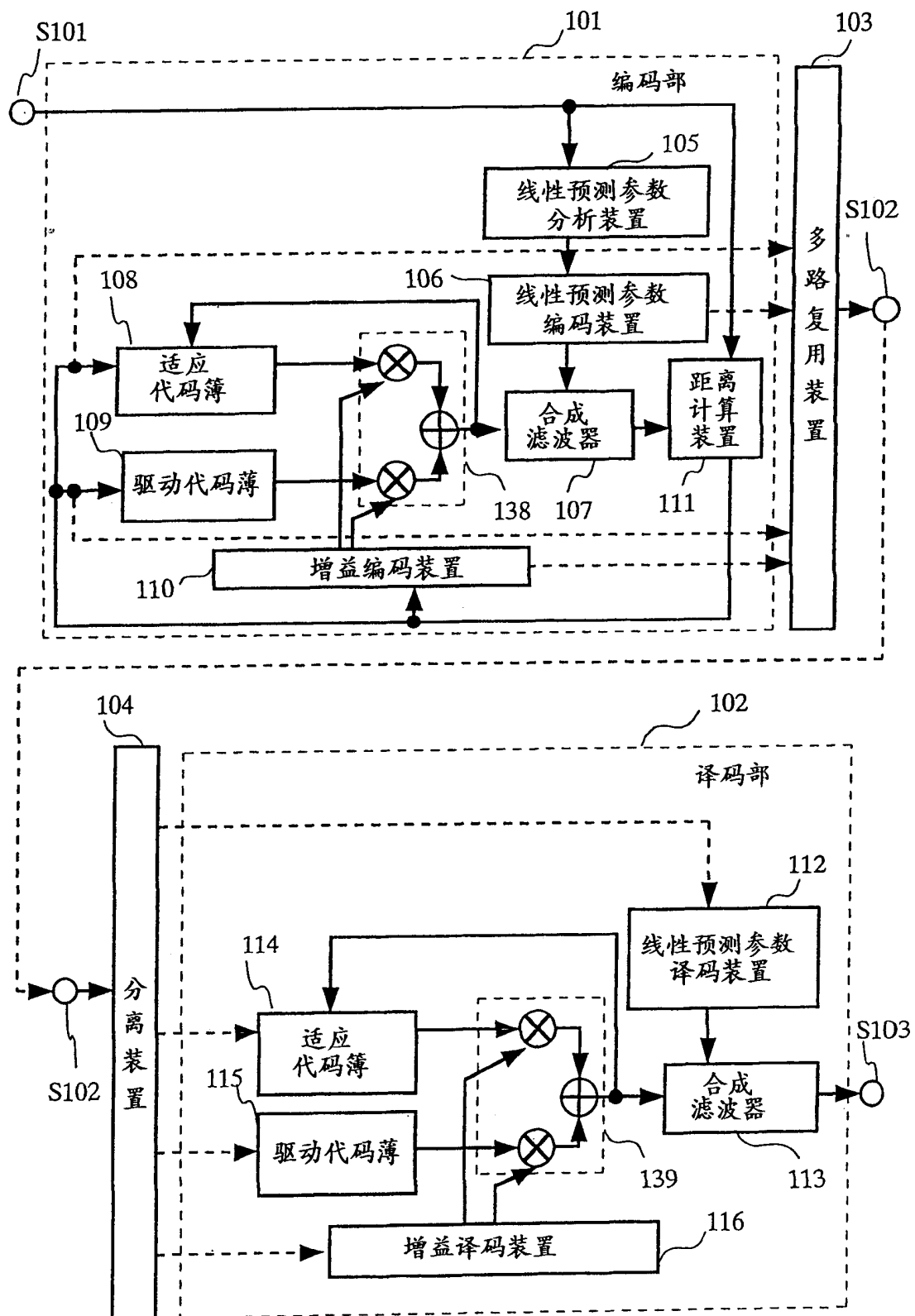


图 6

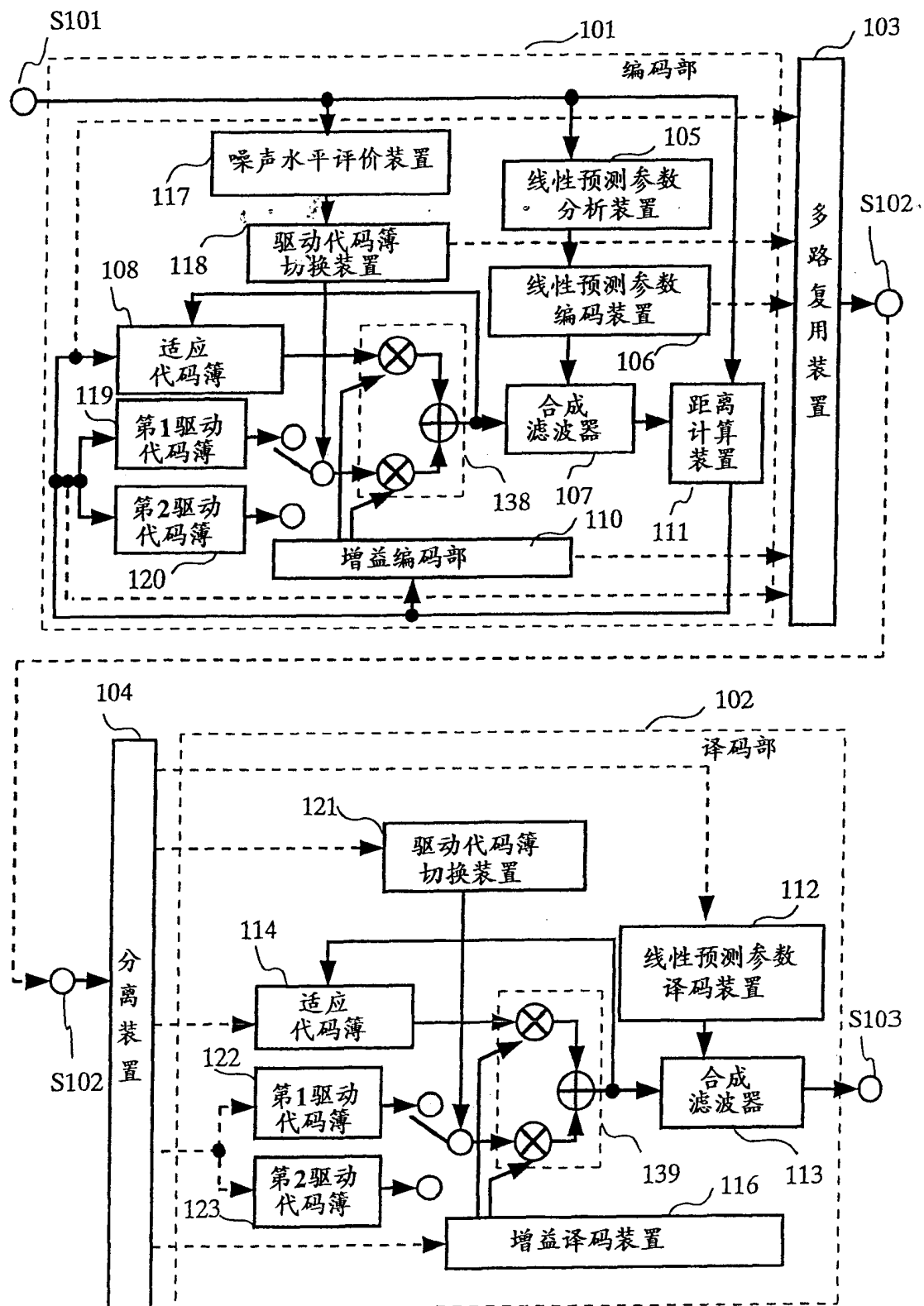


图 7