

混响的混合高斯概率密度建模

卫红凯, 王平波, 蔡志明, 蒋来兴

(海军工程大学电子工程学院, 武汉 430033)

摘要: 混合高斯模型能够有效地拟合混响背景的一维概率密度分布。常用的混合高斯概率密度模型参数估计方法是 EM 迭代算法, 但这种算法的主要缺点是估计精度过分依赖于初始值。而 GreedyEM 算法通过往混合模型中不断地加入高斯分量, 能很好地解决这一问题。文章将多维图象处理中的 GreedyEM 算法加以合理简化, 并给出模型自动定阶方法, 从而成功应用于水声混响的一维混合高斯模型建模中。实验结果表明: 应用新算法能从混响接收数据中准确拟合其概率密度曲线, 并且能适应不同的数据长度, 具有很好的通用性。

关键词: 混合高斯; 概率密度模型; EM; GreedyEM

中图分类号: TJ630

文献标识码: A

文章编号: 1000-3630(2007)-03-0514-05

Gaussian mixture model for reverberation

WEI Hong-kai, WANG Ping-bo, CAI Zhi-ming, JING Lai-xing

(Electronic Engineering College, Navy Engineering University, Wuhan 430033, China)

Abstract: The probability density distribution of one-dimension under the reverberation background can be efficiently modeled using Gaussian mixtures. EM is one of the popular algorithms for parameters estimation of Gaussian mixture probability density model. However, this method highly depends on the initial parameters. The GreedyEM algorithm can solve this problem efficiently by incrementally adding Gaussian components to the mixture. It has been used in multidimensional image manipulations. We properly simplify and apply it to learn Gaussian mixtures of onedimension reverberation successfully, present a method for determining the number of mixing components. We provide experimental results illustrating that the new algorithm can approximate the probability density curve of the receiving reverberation data, and has the potential to be used for different length of data due to its good adaptability.

Key words: Gaussian mixture; probability density model; EM; GreedyEM

1 引 言

经典的主动声纳信号参量检测理论通常假设干扰背景是高斯的, 这种假设下信号的最佳检测就是匹配滤波(或相关检测), 而实际应用中, 如果目标回波位于混响抑制区, 匹配滤波的性能会大为下降。因此, 近些年来, 人们开始研究使用非高斯模型来拟合混响的统计特性。而混合高斯概率密度模型(简记为 GM 模型)具有结构简明、模型参数少、拟合性高、

适用范围广等突出优点, 在图象处理、医疗信号处理、语音信号处理、通信等许多领域被广泛应用作一种典型的非高斯信号概率密度(PDF)模型。本文中也采用 GM 模型来对混响进行概率密度建模。

目前常用的 GM 模型参数估计方法是 EM 迭代算法^[1]。虽然该法保证了在迭代过程中似然函数不会减小, 但它也存在一些不容忽视的缺陷, 比如: 估计精度过分地依赖于初值, 若初值设置不当, 则可能收敛于错误的局部极值点; 无法估计模型阶数, 只能根据经验先验地假定。为了克服 EM 迭代算法的这些缺陷, Verbeek 在图像处理中提出了一种改进的算法-GreedyEM 算法^[2], 它不依赖于初始值, 同时还允许自动地确定 GM 模型的最佳阶数。但文献

中, GreedyEM 算法针对的是图象处理中常用的多维 GM 模型参数估计问题, 结构庞杂, 运算效率低下。把这一算法加以合理简化, 并给出模型自动定阶方法, 成功地应用到水声混响的一维 GM 建模中。对混响实验数据的概率密度拟合结果表明, 基于 GreedyEM 算法得到的 GM 模型参数, 能够很好逼近海洋混响的 PDF。而且, 新算法能适应不同的数据长度, 具有很好的通用性。

2 GM 建模及 EM 迭代算法描述

一般地, GM 模型概率密度(PDF)表述如式(1)所示,

$$f_M(x_n) = \sum_{m=1}^M \varepsilon_m \cdot \varphi_m(x_n) \quad (1)$$

这里, M 为模型阶数。 ε_m 为混合参数(亦称加权系数), 满足式(2)所示的关系。 φ_m 为第 m 个高斯分量的概率密度函数, 表达式如式(3), μ_m 、 σ_m^2 分别为其均值、方差。

$$\sum_{m=1}^M \varepsilon_m = 1 \quad (2)$$

$$\varphi_m(x_n) = \frac{1}{\sqrt{2\pi\sigma_m^2}} \exp\left[-\frac{(x_n - \mu_m)^2}{2\sigma_m^2}\right] \quad (3)$$

可以看到, GM 模型的 PDF 是一种级数形式, 是一个纯粹的数学模型。理论上可以对任何窄带或宽带的单钟型、多钟型概率密度进行拟合。必须指出, 虽然 GM 模型是数学上的拟合, 但并非没有物理意义。事实上, 式(1)能反映多数非高斯源的物理机理, 即非高斯源是由多个高斯源依据一定的概率分布抽选而成。

对于式(1)所示的 GM 模型参数估计问题可描述为: 给定样本数据 $x = \{x_1, x_2, \dots, x_N\}$, 如何确定其概率密度函数所包含的参数组 $\theta = \{\varepsilon_1, \dots, \varepsilon_M; \mu_1, \dots, \mu_M; \sigma_1^2, \dots, \sigma_M^2\}$ 的值, 使以下的对数似然函数 $L_M(\theta)$ 最大。

$$L_M(\theta) = \sum_{n=1}^N \ln f_M(x_n) \quad (4)$$

Aaron 提出 EM 迭代算法^[3], 对事先给定的 GM 模型阶数 M , 用式(5)对参数组 θ 进行更新, 直至 θ 满足一定的终止条件, 此时的对数似然函数即为最大。

x_n 属于第 m 个高斯分量的先验概率 $P(m|x_n)$ 表示为:

$$P(m|x_n) = \frac{\varepsilon_m \varphi_m(x_n)}{f_M(x_n)} \quad (5)$$

则 GM 模型第 m ($1 \leq m \leq M$) 个高斯分量的概率密度函数, 其参数 $\theta_m = (\varepsilon_m, \mu_m, \sigma_m^2)$ 的迭代公式为:

$$\begin{aligned} \varepsilon_m &= \frac{1}{N} \sum_{n=1}^N P(m|x_n) \\ \mu_m &= \frac{\sum_{n=1}^N P(m|x_n) x_n}{\sum_{n=1}^N P(m|x_n)} \\ \sigma_m^2 &= \frac{\sum_{n=1}^N P(m|x_n) (x_n - \mu_m)^2}{\sum_{n=1}^N P(m|x_n)} \end{aligned}$$

3 GreedyEM 算法

GreedyEM 算法通过不断地向 GM 模型中加入高斯分量, 直至使似然函数达到最大来建立 GM 模型, 从而估计出数据的 PDF。

假设已得到 k 阶的 GM 模型概率密度 $f_k(x_n)$, 加入一个高斯分量 $\varphi(x_n, \beta_{k+1})$, 则可得到 $k+1$ 阶的 GM 模型概率密度 $f_{k+1}(x_n)$, 如下式所示:

$$f_{k+1}(x_n) = (1 - \alpha) f_k(x_n) + \alpha \varphi(x_n, \beta_{k+1}) \quad (6)$$

其中, α 是权重系数(或称加权系数), $\alpha \in (0, 1)$, 高斯分量的参数 $\beta_{k+1} = (\mu_{k+1}, \sigma_{k+1}^2)$ 。这样, 对于每一个 k ($1 \leq k \leq M$), 此时可以认为 $f_k(x_n)$ 是确定的, 只要选择合适的权重系数 α 、高斯分量参数 β_{k+1} , 就可以使下式所示的对数似然函数 $L_{k+1}(\theta)$ 达到最大。

$$\begin{aligned} L_{k+1}(\theta) &= \sum_{n=1}^N \ln f_{k+1}(x_n) \\ &= \sum_{n=1}^N \ln [(1 - \alpha) f_k(x_n) + \alpha \varphi(x_n, \beta_{k+1})] \end{aligned} \quad (7)$$

由此可以看出 Greedy EM 算法不像 EM 迭代算法那样, 需要对所有的高斯分量参数都进行初始化, 而是以一个高斯分量 $f_1(x_n)$ 参数的初始化开始, 然后进行如下的迭代:

- (i) 加入一个新的高斯分量;
- (ii) 运用 EM 迭代公式直至收敛到所需精度。

重复以上两步, 直至根据一定的判决准则找到最佳阶数 M 。由于 Greedy EM 算法每次只对两个分量进行处理, 这样就减少了对初始值的依赖。但对于每一次加入的高斯分量 $\varphi(x_n, \beta_{k+1})$, 需要寻找使对数似然函数 $L_{k+1}(\theta)$ 最大的参数 β_{k+1} , 以此作为该高斯分量的参数。

3.1 关于寻找最佳高斯分量的方法

假设已得到 k 阶的 GM 模型概率密度 $f_k(x_n)$, 求 α, β_{k+1} , 使式 (7) 最大。若先认为 $L_{k+1}(\theta)$ 只是 α 的函数, 运用泰勒公式在 $\alpha=0.5$ 处进行二阶展开, 求得使 $L_{k+1}(\theta)$ 最大的 α ; 而 β_{k+1} 中的方差由理论公式给出, 在样本数据 x 中寻找均值使对数似然函数最大^[4]。但这样运算量很大, 尤其当样本数据较大时, 实现起来更显得困难。同时, 因高斯分量方差由理论给定, 使得对数据小概率密度区域的逼近效果差, 下面介绍一种更有效的方法。

根据先验概率密度函数 $P(m|x_n)$, 将给定的样本数据 x 分成 k 个互不相交的子集 $A_i (1 \leq i \leq k)$, 其中, $A_i = \{x_n \in x : P(i|x_n) = \max\{P(m|x_n) | m=1\}^k\}$, 为了提高收敛精度, 可以对各子集 A_i 再进行分块。方法是在 A_i 中均匀地选取两个数 a, b , 把 A_i 分成互不相交的两个子块 A_{i1}, A_{i2} , 其中, A_{i1} 中元素与 a 的距离比与 b 的距离近, A_{i2} 中元素与 b 的距离比与 a 的距离近。各子块的均值和方差作为新加入高斯分量的初始参数, 权重 a_i 设为 $\varepsilon_i/2$, 运用同样的方法, 也可对 A_{i1}, A_{i2} 再进行分块, 直至达到所要求。本文中, 将各 A_i 分成两子块, 收敛精度较好, 符合要求。

有了初始化高斯参数和权重, 用式 (8) 对参量进行更新 (即 partial EM 迭代算法), 选择使对数似然函数最大的一组参量作为新加入高斯分量的参数和权重。式中, A_{ij} 表示对子集 A_i 进行分块产生的子块, 本文中 $j=1, 2$, N_{ij} 表示 A_{ij} 中元素的个数。

$$P(k+1|x_n) = \frac{a_i \varphi(x_n, \beta_{k+1})}{(1-a_i) f_k(x_n) + a_i \varphi(x_n, \beta_{k+1})}$$

$$a_i = \frac{\sum_{x_n \in A_{ij}} P(k+1|x_n)}{N_{ij}}$$

$$\mu_{k+1} = \frac{\sum_{x_n \in A_{ij}} P(k+1|x_n) x_n}{\sum_{x_n \in A_{ij}} P(k+1|x_n)}$$

$$\sigma_m^2 = \frac{\sum_{x_n \in A_{ij}} P(k+1|x_n) (x_n - \mu_{k+1})^2}{\sum_{x_n \in A_{ij}} P(k+1|x_n)} \quad (8)$$

3.2 算法流程

对给定的样本序列 x , GM 模型参数估计算法流程如图 1 所示。

$f_1(x_n)$ 参量的初始化, 就是将样本序列 x 的均值和方差作为 $f_1(x_n)$ 的均值和方差, 加权系数设为 1。式 (5)、式 (8) 的循环终止条件是

$$\sum_{j=1}^{3M} \left| \frac{\theta_j^{(i)} - \theta_j^{(i-1)}}{\theta_j^{(i)}} \right| < R \quad (9)$$

这里, $\theta_j^{(i)}$ 表示第 i 次参数估计向量的第 j 个元素; R 为预先设定的估计改善上限。

随着高斯分量的不断加入, 似然函数会不断变化, 若加入高斯分量后使得似然函数变小, 则停止加入新的高斯分量, 模型阶数随之确定。

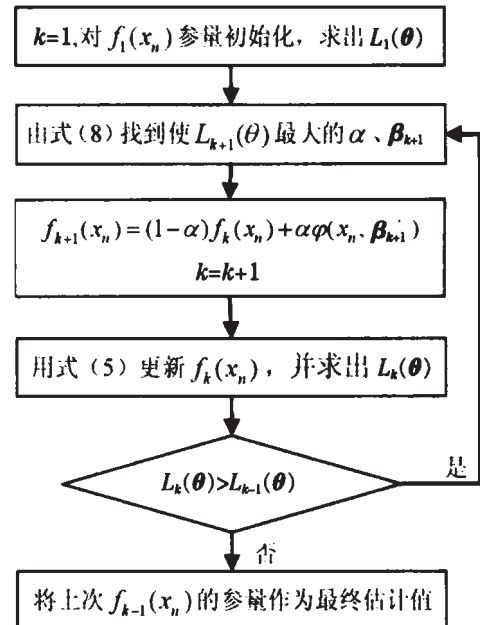


图 1 Greedy EM 估计算法流程图

Fig.1 The flow chart of Greedy EM algorithm

4 混响实验数据建模实例

本文所用的基元级混响实验数据是 2004 年 11 月在海南三亚某区域测量的, 激励信号是 650Hz~850Hz 的线性调频信号, 脉宽为 432ms, 采样频率为 20kHz, 一个周期内的采样点数为 8640。对数据先进行常规波束形成, 然后疏样 4 倍, 经过 FIR 带通滤波后, 得到的波束级数据在一个周期内的采样点数为 2160。

本文用 Greedy EM 算法拟合混响的 PDF 曲线。其中, 图 2 考察了不同点数下, Greedy EM 算法对基元级混响数据的 PDF 曲线拟合能力, 图 3 表征了 Greedy EM 算法对疏样滤波后波束级数据的 PDF 曲线拟合能力。图 2、3 中数据的估计参量如表 1 所示。图 4 通过 Greedy EM 算法与 EM 算法拟合精度的比较, 说明了初始参数对 EM 算法拟合精度的影响, 这两种方法对模型参量的估计结果如表 2 所示。其中, EM 迭代采用文献[5]所用的随机单一初值方案。

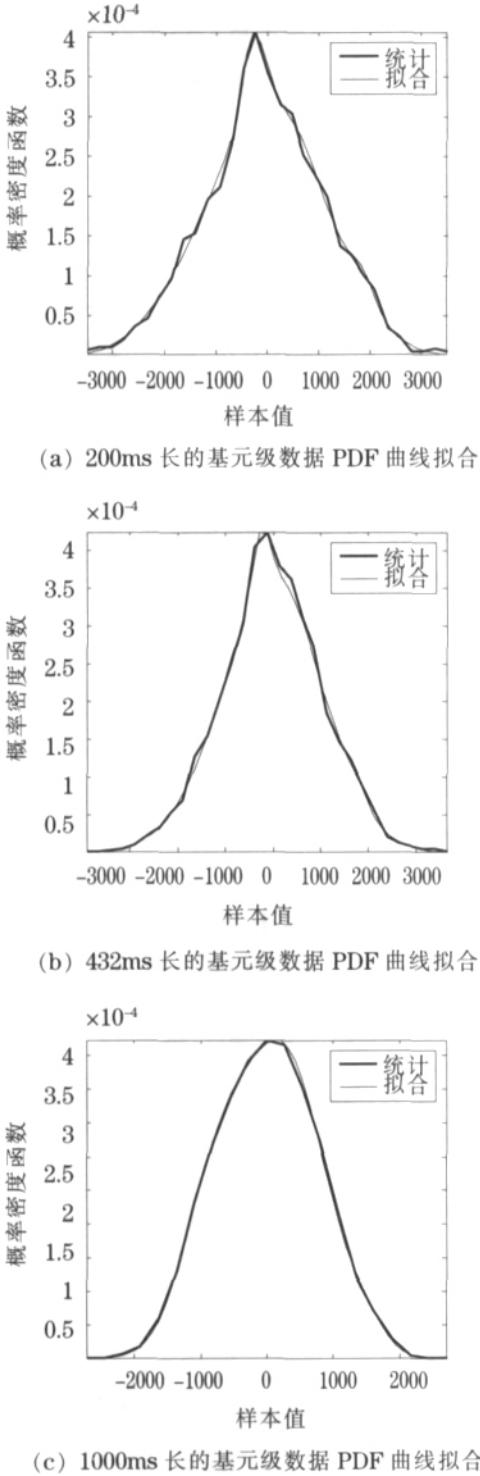


图 2 不同时间长度的基元级数据 PDF 曲线拟合
Fig.2 The PDF curve approximation of original date

为了考察 Greedy EM 算法对小容量数据概率密度曲线的拟合能力,图 2(a)对 200ms 长(4000 点)的基元级数据的 PDF 曲线进行模拟,拟合效果较好,这说明 Greedy EM 算法对可用统计信息量少的小容量数据是适用的;在信号检测时,通常是对一个脉冲宽度的数据进行处理,因此,用图 2(b)来模拟一个脉宽的基元级数据的 PDF 曲线,统计曲线与拟

合曲线的重合程度较高,这说明可以将 Greedy EM 算法应用于混响信号的检测器中,对信号进行 GM 建模及参量估计。图 2(c)表明 Greedy EM 算法对大容量数据的概率密度曲线的拟合精度较高,这是因为数据量大,尽管其非高斯性强,但已知的统计信息多,可以利用的信息就多,也会得到较精确的结果。图 2 中的 PDF 曲线拟合程度都比较高,这说明 GreedyEM 算法不仅参数估计比较准确;而且由于它能适应不同的数据长度,具有很好的通用性。

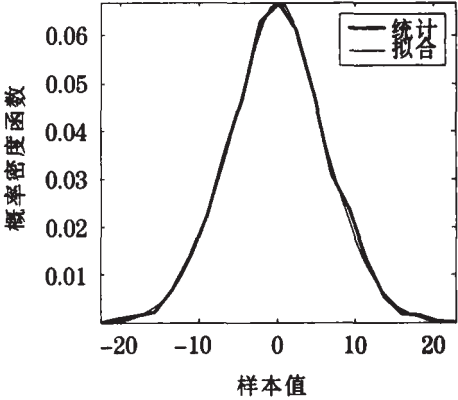


图 3 波束级数据一个脉宽(432ms)的 PDF 曲线拟合
Fig.3 The PDF curve approximation of filtered data for one period (432ms)

为了提高检测性能,通常进入到检测器的信号都是经过滤波处理的。因此,有必要对基元级数据进行疏样滤波处理,考察 Greedy EM 算法对滤波后波束级数据 PDF 曲线的拟合精度。图 3 表征了滤波疏样后 Greedy EM 算法对一个脉宽(2160 点)波束级数据的 PDF 曲线拟合能力,由图可看出拟合曲线能很好地反映出数据的概率分布,准确地模拟出数据的峰值,与统计曲线重合程度很高,比没有经过任何

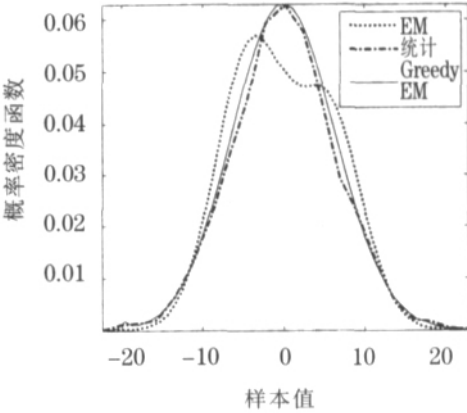


图 4 GreedyEM 与 EM 算法拟合 PDF 曲线比较
(初始参量对后者精度有影响,对前者没有影响)
Fig.4 Comparison of GreedyEM and EM algorithm estimators (initial parameters influence the latter algorithm but do not influence the former)

表 1 图 2、3 中实验数据的 PDF 参量估计值
Table 1 Parameters estimation of the data in figure 2 and figure 3

数据	图 2a	图 2b	图 2c	图 3
ε_1	0.882	0.944	0.701	0.208
μ_1	- 12.479	- 6.842	281.934	1.607
$\sigma^2/10^6$	1.243	1.049	0.402	2.47×10^{-5}
ε_2	1.88×10^{-2}	1.84×10^{-3}	0.283	0.094
μ_2	- 80.947	1762.052	- 776.285	7.789
$\sigma^2/10^6$	3.05×10^{-2}	9.03×10^{-6}	0.246	2.75×10^{-5}
ε_3	3.48×10^{-2}	1.34×10^{-2}	0.0145	0.253
μ_3	- 343.727	- 2052.788	1542.025	4.24×10^{-3}
$\sigma^2/10^6$	2.84×10^{-2}	0.377	0.0491	1.09×10^{-5}
ε_4	2.12×10^{-2}	1.02×10^{-2}	5.29×10^{-4}	0.173
μ_4	1874.655	1806.029	- 3.286	4.178
$\sigma^2/10^6$	6.63×10^{-2}	5.55×10^{-2}	0.115	2.07×10^{-5}
ε_5	0.01	0.03	9.71×10^{-4}	0.128
μ_5	778.804	- 245.288	174.988	- 8.068
$\sigma^2/10^6$	9.14×10^{-2}	2.69×10^{-2}	1.07×10^{-5}	2.35×10^{-5}
ε_6	3.32×10^{-2}	5.6×10^{-4}		0.144
μ_6	- 1945.168	- 1256.78		- 5.228
$\sigma^2/10^6$	0.438	0.134		1.18×10^{-5}

表 2 GreedyEM 与 EM 算法的参量估计值比较
Table 2 Comparison of GreedyEM and EM algorithm parameters estimation

数据	EM算法估计值	GreedyEM算法估计值
ε_1	0.322	0.96
μ_1	- 4.14	0.215
$\sigma^2/10^6$	17.165	39.231
ε_2	0.278	2.1×10^{-3}
μ_2	- 3.568	- 19.633
$\sigma^2/10^6$	21.557	0.233
ε_3	0.4	3.79×10^{-2}
μ_3	5.807	- 4.554
$\sigma^2/10^6$	15.512	28.788

滤波处理的基元级数据的 PDF 曲线拟合精度略高, 这种效果是可以预料的, 因为滤波滤去了基元级数据中的一些干扰频率成分, 使估计的精度变高。

图 4 中传统的 EM 算法拟合出数据的概率密度曲线有两个分离地均值, 这由表 2 也可以看出, 由此产生了两个峰, 与只有一个均值的统计曲线误差较大, 产生了比较严重的偏离, 使拟合性能变差, 这

主要是由于该法对初始值的高度依赖所致; 而 Greedy EM 算法由于不依赖于初始值, 拟合曲线能很好地反映实际统计曲线的特征。这说明, 与 EM 算法相比, Greedy EM 算法有更好的性能。

5 结束语

GM 模型是对非高斯数据进行 PDF 建模的最有效模型之一。尽管 EM 算法是 GM 建模的最常用的算法, 但其存在需要先验地假定 GM 模型阶数、精度对初始值过分依赖等缺点, 而 Greedy EM 算法能有效克服 EM 算法的这些缺点。本文将 Greedy EM 算法进行简化与定阶处理, 成功应用到混响的 GM 建模及参量估计中。实验结果表明: 该法对混响 PDF 曲线的拟合性能比 EM 迭代算法的好, 并且能适应不同的数据长度, 有很好的通用性。这在非高斯信号检测处理领域中有着重要的实际意义。

通常, 传统的 GM 建模方法(如 EM 迭代算法), 在处理有色非高斯数据时, 需要先对数据进行有效的预白化处理, 才能取得较好的建模性能, 而本文用 GreedyEM 算法对基元级(没有经过任何处理, 有色非高斯)混响实验数据的 PDF 曲线的拟合结果表明: 不需要进行预白化处理, 也可得到较精确的结果。

参 考 文 献

[1] Dempster A P, Laird N M, Rubin D B. Maximum likelihood from incomplete data via the EM algorithm[J]. Roy Statist Soc, 1977, 39: 1-38.

[2] Verbeek J J, Vlassis N, Krose B. Efficient greedy learning of gaussian mixture models[R]. The Netherlands: Computer Science Institute, University of Amsterdam, 2001.

[3] Render R A, Walker H F. Mixture densities, maximum likelihood and the EM algorithm[J]. SIAM Review, 1984, 26(2): 195-239.

[4] Nikos Vlassis Aristidis Likas. A greedy EM algorithm for Gaussian mixture learning[J]. Neural Processing Letters, 2000, 15(1): 77-87.

[5] Shawn M Verbout, James M Ooi, Jeffrey T Ludwig, Alan V Oppenheim. Parameter estimation for autoregressive Gaussian-mixture processes: the EMAX algorithm[J]. IEEE Transactions on Signal Processing, 1998, 46(10): 2744-2756.