

贝叶斯概率图像分割中混合模型参数高效计算的研究

Study on Effective Computation of Mixture Model Parameters for Bayesian Probabilistic Image Automatic Segment

郭 平

(北京师范大学信息科学学院计算机系 北京100875)

(中科院自动化所模式识别国家重点实验室 北京100080)

Abstract When image segmentation is treated as a problem of clustering pixels, statistical finite mixture model can be used to classify image samples. With estimated mixture model parameters, we can use some Information Theoretical Criteria to determine how many regions should be segmented on a given image without a priori knowledge to conduct automatic image segmentation. In this paper, we consider the problem in practical implementing segmentation based on mixture models and suggest combining several techniques such as data reduction, Competitive Learning and variant EM algorithm to effectively estimate mixture parameters and reduce intensive computation task. The combined technique is of significance for automatic image segmentation in real-time or near real-time applications.

Keywords Automatic image segmentation, Region number determination, Finite mixture model, Data reduction, EM algorithm

1. 引言

图像分割的目的是将图像划分为一些互不重叠的区域,这是计算机视觉中的一个重要研究领域,也是图像理解的基础。在众多的图像分割技术中^[1],特征空间聚类可以说是最常用的方法之一。通常用一确定的特征表征同一分割区域的像素。这些特征被量化成特征变量,同一分割区域的像素的特征变量基本上有类似的数值,不同分割区域的像素特征变量数值不同。在实施图像分割时,首先在特征空间把特征变量聚类,然后把特征空间的每一点映射回到图像空间的像素。

当把图像分割看成是聚类过程时,有多种技术可以应用于聚类。例如众所周知的k-均值算法,竞争学习法(CL)等^[2,3],而用最大似然近似方法去匹配混合模型则是用得最多的方法^[4,5]。由于混合模型分析具有多种优点,在近几年把混合模型分析应用于图像分割已经吸引了一些研究者的兴趣^[6,7]。

把有限混合模型分布应用于图像分割是基于一个假定,即认为特征空间中的数据是服从k个概率密度混合分布的样本。这是一种基于统计模型的方法,k在这里可以看成是类别数目。如果能正确选取k,可以得到较好的图像分割结果。否则,样本不能被适当聚类,图像分割效果也不可能好。而k的确定是机器学习中最困难的事情之一。过去一般的做法是预先假定一个k,然后进行图像分割。这种做法影响了由机器进行全自动图像分割的能力。

若假定要分割的图像区域数目与特征空间的类别数目是相等的,则可以用信息理论准则去判定一幅图像应该分割成多少区域。在确定要分割的图像区域数目后并用贝叶斯概率分割图像,这样可以实现机器全自动图像分割^[8,9]。用贝叶斯概率分割图像,前提是首先要基于图像样本来估计混合模型参数。目前普遍采用的算法是用EM算法来估计混合模型参数。虽然用混合模型分布结合EM算法来进行图像分割有其优势,但是有个明显的不利因素是计算工作量较大,这也是基于统计模型分析的通常的弱点,给一些需要在线图像分割的

实际问题带来困难。本文提出了采用组合多种技术方法,象数据简约、竞争学习和变异EM算法等方法,去高效地计算混合模型参数,提高计算速度,为实现实时图分割自动分割提供一定的基础。

2. 相关背景

2.1 有限混合模型聚类

对一幅彩色图像来说,特征空间可以选择为RGB或其他色彩空间。如果给定的图像有N个像素,用 x_i 表示在第i个像素的观测值。图像中所有的样本构成数据组 $D=\{x_i\}_{i=1}^N$ 。在假定特征空间中的数据点是服从k个高斯概率密度混合分布的样本后,其联合概率分布可以写成:

$$p(x, \Theta) = \sum_{j=1}^k \alpha_j G(x, m_j, \Sigma_j) \quad (1)$$

$$\text{with } \alpha_j \geq 0, \sum_{j=1}^k \alpha_j = 1$$

其中 $G(x, m_j, \Sigma_j) = \frac{1}{(2\pi)^{d/2} |\Sigma_j|^{1/2}} \times \text{Exp} \left[-\frac{1}{2} (x - m_j)^T \Sigma_j^{-1} (x - m_j) \right]$

是多变量高斯概率密度分布的一般表示。式中 x 代表随机矢量, d 是随机矢量 x 的维数,而 $\Theta = \{\alpha_j, m_j, \Sigma_j\}_{j=1}^k$ 是系列混合模型参数。对混合模型的第j个分量而言, α_j 是混合模型的第j个分量的混合权重, m_j 是均值矢量,而 Σ_j 是协方差矩阵。这些参数我们预先并不知道,用多少个高斯分量可以最佳匹配样本概率分布,也无法预先得知。通常只能假定一个k值,根据给定的图像样本用最大似然估计,利用EM算法来估计它们^[4,5]。EM算法如下:

E-步骤: 计算后验概率

$$p(j|x_i) = \frac{\alpha_j G(x_i, m_j, \Sigma_j)}{\sum_{l=1}^k \alpha_l G(x_i, m_l, \Sigma_l)} \quad (2)$$

M-步骤: 计算混合模型参数

$$\alpha_j^{\text{new}} = \frac{1}{N} \sum_{i=1}^N \frac{\alpha_j G(x_i, m_j, \Sigma_j)}{\sum_{l=1}^k \alpha_l G(x_i, m_l, \Sigma_l)}$$

$$\begin{aligned}
&= \frac{1}{N} \sum_{j=1}^N p(j|x_i) \\
m_j^{nw} &= \frac{1}{a_j N} \sum_{i=1}^N p(j|x_i) x_i \\
\sum_j m_j^{nw} &= \frac{\sum_{j=1}^N p(j|x_i) (x_i - m_j) (x_i - m_j)^T}{\sum_{j=1}^N p(j|x_i)}
\end{aligned} \quad (3)$$

这 E-M 两步交替迭代计算,直到系统的似然函数收敛到某一局部极小点。

2.2 模型选择准则

当估计到混合模型参数值后,可以应用信息理论准则去确定要分割的区域数目。在文献中有一些信息理论准则可以用于模型选择,例如, AIC^[10], AICB^[11], CAIC^[12], SIC^[13]等,这些准则把系统似然函数的极大值与达到该值的参数数目结合起来去确定模型选择。在本文中我们采用 SIC 来确定要分割的区域数目。

$$SIC = J_2(k) = -2\ln[\max[L(k)]] + (2 + \ln N)m_k$$

$$m_k = kd(k-1) + kd(d+1)/2 \quad (4)$$

式中 $L(k)$ 是系统在由 k 个模型构成时的似然函数, m_k 是惩罚项。这些准则都比较类似,不同之处在于惩罚项不同。通常的做法是先从 $k=1$ 开始,估计混合模型参数,计算 $J_2(k=1)$ 值,然后令 $k=2$,再计算 $J_2(k=2)$ 。对每个 k 下估计混合模型参数,计算函数 $J_2(k)$ 。当得到一系列的 $J_2(k)$ 后,选最小的函数值 $J_2(k)$ 对应的 k^* ,即所要分割的区域数目 k^* 为:

$$k^* = \operatorname{argmin}_k (J_2(k)) \quad (5)$$

2.3 贝叶斯概率分类

在确定 k 后,我们也同时得到了在这个 k 下的后验概率函数(式2)。后验概率公式表示样本点 x 属于类别 j 的概率。对给定的 x_i 样本点,可以得到 k 个概率 $p = (j=1|x_i), p(j=2|x_i), \dots, p(j=k|x_i)$, 利用贝叶斯定则, $j^* = \operatorname{argmax}_j p(j|x_i)$, 把样本 x_i 划分到类别 j^* 中,这个过程称为贝叶斯概率分类。因此我们可以看出,有限混合模型图像分割实际上是一种贝叶斯概率分类像元的过程。

3. 简约数据与算法

从上可见,要实现图像自动分割,首先要基于图像特征空间的样本来估计混合模型参数。

在实际图像处理中,即使对一幅大小为 128×128 像元的彩色图像,在特征空间中也有 $N=16384$ 个样本点,基于所有样本用 EM 算法去估计密度分布的计算强度是很大的。要解决该问题,可采用的方法之一是从训练样本中提取一代表性的子集,用该子集去估计参数,以降低计算强度。实际上使用简约的样本子集来代表一很大的样本集的概念由来已久,并且应用范围广泛,例如在众所周知的 k -均值算法,竞争学习方法等算法中。

使用简约的样本子集去构造贝叶斯概率分类器有多种优点。首先因为使用的样本数目 N_s 远远小于 N ,这就极大地节省了计算时间。其次抽样时可以提取图像的主要特征且忽略次要图像区域。另外抽样也可以稀疏样本并在某种程度上减低各类的重叠,这使聚类过程变得相对容易。

得到简约样本集的方法是从原始样本集中抽取样本点。抽样的目标是,既减小了样本集的尺度又保持了与原始样本集类似的统计性质。目前有好几种选择样本的方法,例如简单随机抽样法(simple random sampling, SRS), 分层抽样法等^[15,16]。在实际抽样时,分层抽样法的效果可能比较好,但它

要求知道样本集的先验信息。不幸的是我们在图像自动机器处理中,并没有给定图像的先验信息。因此在本研究工作中我们认为 SRS 才是一种切实可行的有效技术。

SRS 是一种从 N 个样本中抽取 N_s 个样本,在 N 中的每一样本点都有相等概率被选取的技术。在应用 SRS 时,每一样本点被选取的概率是 $1/N$ 。由于概率相等,应用 SRS 技术得到的子集中各类的比率不变。因此简约样本集的尺度正比于原始样本集的尺度,原始样本集中的某类的尺度大,简约样本集中的该类的尺度也大。

原理上,训练样本的简约子集应该用如下所述的方式选择:用简约子集估计的密度函数应与原始样本集估计的函数尽可能地匹配,使得在两种估计下选择的分割区域数目一致,而且具有最小的分类误差。在使用 SRS 技术时,问题转化为如何选择合适的子集样本数目 N_s ,但在实际中选择合适的 N_s 是非常困难的事情。如我们所知,用太小的有限样本尺度,不可能得到较为精确的各个类别的参数估计。因此, N_s 应按最小容忍误差原则选择。如果知道一些先验知识,在 Parzen 密度估计时,可根据 Fukunaga 提出的算法来选取较好的子集^[17]。但是实际情况我们没有任何先验知识,唯一方法是采用试错法(trial-and-error)来确定 N_s ,这种方法又使计算强度变大。我们的目的是考虑如何进行高效计算,因此本研究工作中 N_s 是根据经验方法设定的。

在抽样得到 N_s 个样本数目后,可采用 EM 算法来估计混合模型参数。常规的 EM 算法是一种迭代算法,在实际情况下其收敛速度很慢。为加速收敛,这里可以采用适当地近似,即采用变异的 EM 算法^[14]来估计混合模型参数。变异的 EM 算法是在对公式(3)中 $p(j|x)$ 进行估计时将其二值化:

$$\begin{aligned}
p(j|x_i) &= I(j|x_i) \\
&= \begin{cases} 1 & \text{if } j = \operatorname{argmax}_j a_j G(x_i, m_j, \Sigma_j) \\ 0 & \text{otherwise} \end{cases} \quad (6)
\end{aligned}$$

采用这种变异的 EM 算法,可显著地加速收敛速度,提高计算效率。但这种方法估计的参数精度对均值的初始化值依赖性很强。我们知道,在迭代算法中混合模型参数的初始化值不仅对收敛速度有较大的影响,而且如果选择的初始参数不佳,就会收敛到所不希望的局部极小点,用变异的 EM 算法更是如此。由于没有先验知识,参数初始化的方法通常采用随机取值法。为解决收敛到不希望的局部极小问题,我们采用固定初始参数值而不是随机取值,即将混合模型的初始均值和协方差值偏置到特定的值。实际的初始均值是用 CL 算法预聚类确定的。采用 CL 算法预聚类得到初始的均值矢量,还能在使用变异的 EM 算法时加速收敛到恰当的局部极小。

4. 实验结果讨论

在实验中,我们使用了一幅合成图像以及“房屋”与“帆船”等一些常用图像,去检验所提出的组合技术的效果。所用合成图像为24位真彩色,大小为 128×128 像元。

在实验中, N_s 选定在2000左右。我们运行变异的 EM 算法估计混合模型参数并计算 $J_2(k)$ 函数。采用 CL 算法初始化的均值矢量,确实可消除 EM 算法收敛到不同的局部极小的影响。对方差矩阵值初始化时,由于我们知道在 RGB 空间颜色值范围是0-255,因此协方差矩阵对角线上的值设为10左右。如果此值太小则没有利用到 CL 算法初始化的优势,还可能使收敛速度变慢。此值太大则会覆盖不止一个类簇,并导致收敛到不希望的局部极小点。

实验结果表明,对合成图像,由于各类簇之间分离的较好,因此采用简约子集对估计区域数目及分割质量并没有什么影响。图1显示了合成图像的简约子集分布。对“房屋”图像,采用简约子集后,原始特征空间的三个主要区域变的比较突出,而原先的一些细节变成了外点或可以看成噪声,对参数估

计的贡献则可忽略不记。在计算效率上,实验结果表明我们所提出的组合技术改进速度是非常显著的。若采用通常的 EM 算法,相继的似然函数误差收敛到 $|L_{i+1}(k, \theta_{i+1}) - L_i(k, \theta_i)| < 10^{-5}$ 需要的迭代次数在200次以上,而采用组合技术,大多数情况下小于10次迭代就可达到。

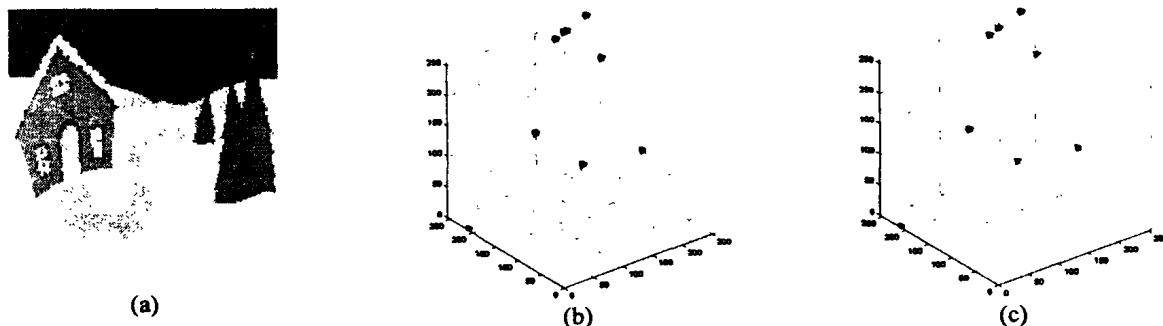


图1 合成图像及 RGB 特征空间:(a)原始图像;(b)特征空间数据分布;(c)简约的数据空间分布

小结 在本文中,我们提出了采用组合数据简约、竞争学习和变异 EM 算法等方法,去高效估计混合模型参数。然后基于 SIC 准则确定图像分割区域数目,用贝叶斯规则把像元分类到适当的区域。采用 SRS 技术获取简约子集,CL 算法初始化均值矢量初值,变异 EM 算法估计混合模型的参数。实验结果表明对统计混合模型图像处理而言,所提出的组合技术是一种很好的高效计算的近似方法。采用简约子集没有明显降低分割质量,但却显著地节省了计算时间。因此,采用组合技术可以高效计算混合模型参数,提高计算速度,为实现实时或近实时图分割自动分割提供一定的基础。

参考文献

- 1 Fu K S, Mui J K. A survey on image segmentation. Pattern Recognition, 1981, 13: 3~16
- 2 Ahalt S C, et al. Competitive learning algorithms for vector quantization. Neural Networks, 1990, 3: 277~290
- 3 Scheunders P, Backer S D. High-dimensional clustering using frequency sensitive competitive learning. Pattern Recognition, 1999, 32: 193~202
- 4 Dempster A P, Laird N M, Rubin D B. Maximum-likelihood from incomplete data via the EM algorithm. J. Royal Statist. Society, 1977, B39: 1~38
- 5 Redner R A, Walker H F. Mixture densities, maximum likelihood and the EM algorithm. SIAM Review, 1984, 26: 195~239
- 6 Santago P, Gage H D. Statistical models of partial volume effect. IEEE Trans. Image Processing, 1995, 4(11): 1531~1540
- 7 Sanjay-Gopal S, Hebert T J. Bayesian Pixel Classification Using Spatially Variant Finite Mixtures and the Generalized EM algorithm. IEEE trans. Image Processing, 1998, 7(7): 1014~1028
- 8 Guo P, Cheung C C, Xu L. Region Number Determination in Automatic Image Segmentation Based on BKYY Model Selection Criterion. In: Proc. of 1999 IEEE-EURASIP Workshop on Nonlinear Signal and Image Processing, 1999. 743~746
- 9 Guo P, Lyu M R. A Study on Color Space Selection for Determining Image Segmentation Region Number. In: Proc. of the 2000 Intl. Conf. on Artificial Intelligence, 2000. 1127~1132
- 10 Akaike H. A New Look at the Statistical Model Identification. IEEE Transactions on Automatic Control, 1974, AC-19: 716~723
- 11 Bozdogan H. Multiple Sample Cluster Analysis and Approaches to Validity Studies in Clustering Individuals. Doctoral Dissertation, University of Illinois at Chicago Circle, 1981
- 12 Bozdogan H. Model Selection and Akaike's Information Criterion: The General Theory and its Analytical Extensions. Psychometrika, 1987, 52(3): 345~370
- 13 Schwarz G. Estimating the Dimension of a Model. The annals of Statistics, 1978, 6(2): 461~464
- 14 Wang T J. Personal communication
- 15 Govindarajulu Z. Elements of sampling theory and methods. Prentice-Hall, Inc. (Upper Saddle River, NJ), 1999
- 16 Colmenares G L, Perez R. A data reduction method to train, test, and validate neural networks. In: Proc. of IEEE Southeastcon '98, 1998. 277~280
- 17 Fukunga K, Hayes R R. The reduced parzen classifier. IEEE Trans. on Pattern Analysis and machine Intelligence, 1989, 11(4): 423~425