

DOI: 10.7652/xjtuxb201308012

手写维文字符分割中的多信息融合路径寻优方法

许亚美, 卢朝阳, 李静, 姚超

(西安电子科技大学综合业务网理论与关键技术国家重点实验室, 710071, 西安)

摘要: 针对维吾尔词书写粘连和手写笔画漂移等问题, 提出一种基于多信息融合路径寻优的字符分割算法。利用笔画提取、切分和聚类, 过分割单词图像得到主体和附加字段, 通过字段模糊匹配获得鲁棒的字根序列描述, 以抑制笔画漂移造成的干扰; 由建立的匹配位置高斯模型来估算字段匹配信息, 经对单字分类器输出进行置信度转换, 从而得到字符识别信息, 再运用数据统计获取单词语义信息; 由构建的字符序列二阶 Markov 语言模型, 基于 Bayes 准则, 提出了单词后验概率的多信息加权融合计算方法, 通过字段匹配及字根合并的路径寻优, 可得到最佳字符分割结果。在手写维文样本库上的实验表明, 所提算法能有效提升字符分割的准确率和稳定性。

关键词: 信息处理技术; 手写文字识别; 字符分割; 维吾尔语; 多信息融合

中图分类号: TP391.4 **文献标志码:** A **文章编号:** 0253-987X(2013)08-0068-06

A Path Optimization Method Based on Multiple Information Fusion for Handwritten Uyghur Character Segmentation

XU Yamei, LU Zhaoyang, LI Jing, YAO Chao

(State Key Laboratory of Integrated Service Networks, Xidian University, Xi'an 710071, China)

Abstract: Character segmentation is a key technique for Uyghur handwriting recognition, but cursive characters and the phenomenon of stroke drift make the segmentation difficult. A new character segmentation algorithm based on multiple information fusion is proposed to solve the problem. Strokes of a word are extracted, segmented and clustered to get two types of sections: main and affix. The robust over-segmentation primitive sequences are obtained using fuzzy section matching to reduce the interference from stroke drift. Then, the matching information is estimated by constructing a matching position Gaussian model. The recognition confidence is converted from character classifier outputs by confidence transformation, and the semantic information is obtained by word data statistics. A character sequences Markov model is presented and the formula to calculate the posterior probability of a word is derived based on the Bayes criterion. The optimal path and the optimal segmentation result are achieved by weighted multiple information fusion. Experiments show that the proposed algorithm can effectively improve the accuracy and stability of character segmentation.

Keywords: information processing technology; handwriting recognition; character segmentation; Uyghur language; multiple information fusion

收稿日期: 2012-11-03。 作者简介: 许亚美(1978—), 女, 博士生; 卢朝阳(通信作者), 男, 教授, 博士生导师。 基金项目: 国家自然科学基金资助项目(60872141); 中央高校基本科研业务费专项资金资助项目(K50510010007); 华为科技基金资助项目(HITC2011023)。

网络出版时间: 2013-05-21

网络出版地址: <http://www.cnki.net/kcms/detail/61.1069.T.20130521.1432.001.html>

维吾尔文(简称维文)是我国重要的少数民族文字,维文的研究有益于促进民族地区信息和科技的发展。目前,手写文字识别的瓶颈是断裂字符的分割问题,词是维文独立运用的最小语言单位^[1],因此词至字符的分割技术成为手写维文识别的关键。维文由多个字符沿着一条想象中的水平轴线(基线)自右至左书写而成,如图 1 所示,其结构规则^[2]为:①词中沿基线书写的部分为主体笔画,其余点、元音符号等为附加笔画;②词由多个连体段组成,各连体段包含一个或多个粘连书写的字符;③书写时粘连字符在主体笔画上相连;④各字符宽度不等;⑤手写维文主体和附加笔画的相对位置较为随意,易产生笔画漂移现象。上述特性也是造成手写维文字符分割的主要难点。

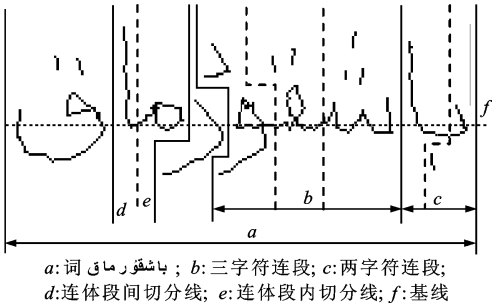


图 1 维文单词结构规则

维文字符分割的研究刚刚起步,相关文献很少,主要集中在印刷体文字分割^[2-3]方面。考虑到维文中部分文字借用了阿拉伯字母,而阿拉伯字符分割技术^[4-10]大体包括 3 类:①基于图像分析的分割方法^[5-6];②基于字符识别的分割方法^[7];③先过分割,再进行合并路径寻优的混合型方法^[8-10]。研究发现,综合多种切分信息是提高路径寻优准确率的有效途径,如:Al Hamad 等提出了综合切分点中心、中间字段和右端字段的识别信息来合并过分割单元的方法,并取得了较好的分割结果^[9];Boukharouba 等将隐马尔科夫模型引入到了字符序列的语义信息中^[10],但未考虑阿文或维文的笔画漂移问题,因而影响了分割效果。

综上,本文针对手写维吾尔词书写时的粘连、交错、笔画漂移等特点,提出了一种基于过分割、模糊匹配和切分路径寻优的字符分割算法。

1 过分割

本文算法的整体框架如图 2 所示,其中涉及到的过分割是将单词图像切分为单个字符或字符组成

部分。

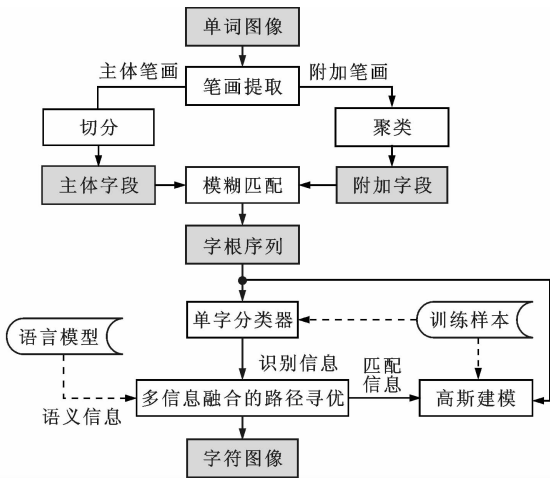


图 2 本文算法的整体框架

现有过分割技术^[2-10]一般是在切分点处直接分割单词图像。手写维文的附加笔画较多且书写位置随意,直接分割很可能将同一字符的附加笔画错分至 2 个字符。考虑到字符在主体笔画粘连的因素,本文提出先通过基线域检测,再结合连通域分析,来提取单词的主体笔画和附加笔画,然后对主体笔画经过切分得到主体字段,同时对附加笔画聚类形成附加字段。

本文过分割算法流程如图 3 所示,步骤如下。

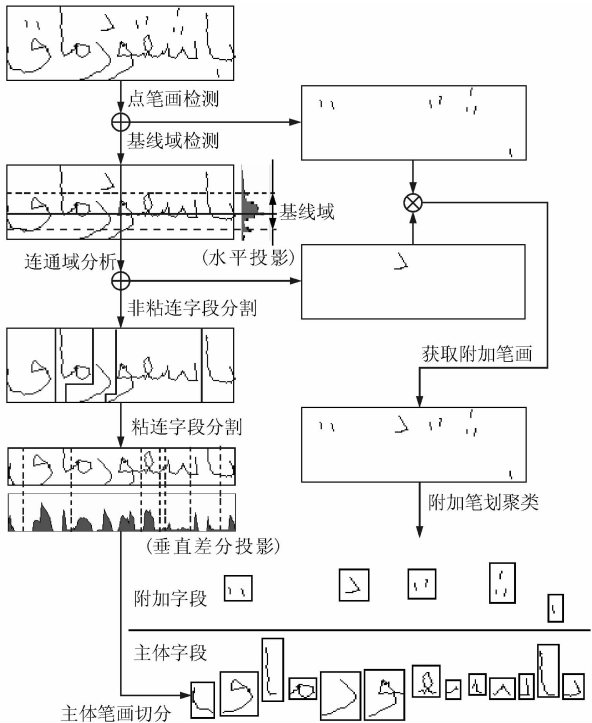


图 3 过分割算法流程

(1)预处理:包括二值化、倾斜校正、细化、归一

化和断笔连接。

(2)点笔画检测:设单词笔画为 $S=(S_1, S_2, \dots, S_n)$, 预设点阈值为 T , 当笔画连通域面积 $\zeta(S_i) < T$ 时, 判定该笔画为点笔画, 全部点笔画记作 $P=(P_1, P_2, \dots, P_l)$ 。

(3)基线域检测:利用剩余笔画 $S_R=S-P$ 检测基线 B 和基线域 $[B_u, B_d]$, B_u 和 B_d 为基线域的上界和下界, 计算式如下

$$\left. \begin{aligned} B &= \arg[\max_j H(j)] \\ B_u &= \arg\left[\sum_{j=k}^B H(j) = \sigma \sum_{j=0}^B H(j)\right] \\ B_d &= \arg\left[\sum_{j=B}^k H(j) = \sigma \sum_{j=B}^{L-1} H(j)\right] \end{aligned} \right\} \quad (1)$$

$$H(j) = \sum_{i=0}^{W-1} S_R(i, j), \quad j = 0, \dots, L-1 \quad (2)$$

式中: $H(j)$ 表示水平投影; W 和 L 为 S_R 的宽和高; σ 决定基线域的大小, 是经验值, 一般 $\sigma=0.7$ 。

(4)连通域分析:对 S_R 中任意两笔画 S_i 和 S_j , 计算其外接矩形框, 如果两矩形框在宽度上的重叠区域大于其中较小宽度的 $2/3$, 那么认为 S_i 和 S_j 是交叠的, 检测其中的离基线相对较远的笔画(假设为 S_i), 若 S_i 不在基线域内, 则将 S_i 归入附加笔画 S_A 。

(5)主体笔画切分:在基线域内对主体笔画 $S_M=S_R-S_A$ 计算垂直差分投影^[9], 取极小值点作为切分点进行分割, 得到主体字段 $M=(M_1, M_2, \dots, M_N)$ 。

(6)附加笔画聚类:采用顺序聚类最大最小算法^[11]对点笔画 P 进行聚类, 根据同字符中点笔画写在基线一侧的规则^[2], 定义点 P_i 和 P_j 的距离测度, 计算式如下

$$D(P_i, P_j) = \begin{cases} +\infty, & [Y(P_i) - B][Y(P_j) - B] < 0 \\ |X(P_i) - X(P_j)|, & \text{其他} \end{cases} \quad (3)$$

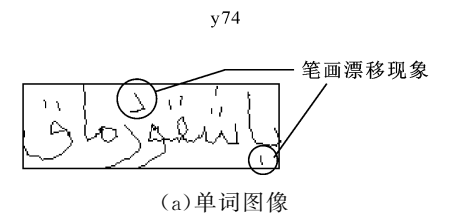
式中: $X(\cdot)$ 和 $Y(\cdot)$ 分别表示点中心的 X 和 Y 轴坐标; B 为基线位置。聚类顺序取 X 轴向自右至左, 聚类后的点群作为整体, 联合上述步骤(4)得到附加笔画 S_A , 由此构成附加字段 $A=(A_1, A_2, \dots, A_M)$ 。

2 切分路径寻优

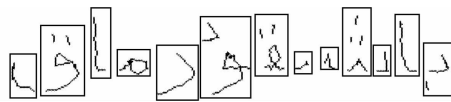
2.1 字段模糊匹配

字段模糊匹配如图4所示。维文中词的书写是先粘连主体笔画, 再添加附加笔画, 因而极易产生笔画漂移, 即某些字符的附加笔画更接近相邻字符(见

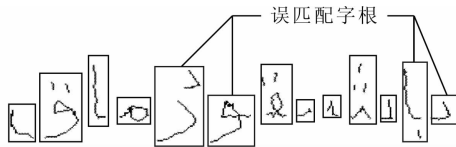
图4a), 若字段以最近距离原则匹配, 则易产生错误字根(见图4c)。为解决这一问题, 研究者采用直接将附加笔画当作一个整体来提取特征^[12], 或者在提取滑动窗特征时通过改变窗体角度来兼顾漂移笔画^[13], 但上述方法均基于整词识别, 不能应用于字符分割上。为此, 本文提出将每个附加字段与其最近距离及左、右相邻的3个主体字段进行模糊匹配(见图4d), 由此获得多个字根图像序列, 然后以匹配置信度对模糊匹配的代价进行反馈, 最后通过路径寻优获得最佳匹配结果。



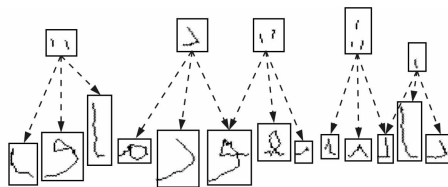
(a) 单词图像



(b) 正确匹配的字根序列



(c) 最近距离原则匹配的字根序列



(d) 模糊匹配关系

图4 字段模糊匹配示意图

考虑到维文中字符由1个主体字段和0~2个附加字段组成^[2], 字符主体与附加字段间的距离

$$d(C_i) = \begin{cases} 0, & C_i = M_i \\ |X(M_i) - X(A_j)|, & C_i = M_i \cup A_j \\ \left| X(M_i) - \frac{X(A_j) + X(A_k)}{2} \right|, & C_i = M_i \cup A_j \cup A_k \end{cases} \quad (4)$$

式中: M_i 为字符 C_i 的主体字段; A_j 或 A_k 为附加字段; $X(\cdot)$ 表示字段外接矩形框中心点的 X 轴坐标。

设 $d(C_i)$ 为独立抽样并满足高斯分布, 于是其

均值 μ_d 和方差 σ_d 可根据最大似然准则估计,即

$$\mu_d = \frac{1}{N} \sum_{i=1}^N d(C_i)$$

$$\sigma_d = \left(\frac{1}{N-1} \sum_{i=1}^N [d(C_i) - \mu_d]^2 \right)^{1/2} \quad (5)$$

式中: N 为训练字符样本数。由此,字符 C_i 的匹配置信度

$$p(G(C_i)) = \frac{1}{\sigma_d(2\pi)^{1/2}} \exp \left\{ -\frac{[d(C_i) - \mu_d]^2}{2\sigma_d^2} \right\} \quad (6)$$

2.2 字符序列 Markov 模型

设单词中各字符的识别结果仅与字符自身图像以及前一个字符的识别结果有关,由此可建立字符序列二阶 Markov 语言模型。设字符图像序列 $\mathbf{C} = (C_1, \dots, C_n)$ 在单字符分类器下的识别结果为 $\mathbf{I} = (I_1, \dots, I_n)$, 利用条件概率公式进行推导,根据字符识别仅与其自身图像及前一字符识别结果相关的假设,则有 $p(I_i | C_1, \dots, C_n) = p(I_i | C_i)$, $p(I_i | I_{i-1}, \dots, I_1, C_1, \dots, C_n) = p(I_i | I_{i-1}, C_i)$, 由此得到

$$p(I_1, \dots, I_n | C_1, \dots, C_n) =$$

$$p(I_1 | C_1, \dots, C_n) \prod_{i=2}^n p(I_i | I_{i-1}, \dots, I_1, C_1, \dots, C_n) =$$

$$p(I_1 | C_1) \prod_{i=2}^n p(I_i | I_{i-1}, C_i) \quad (7)$$

利用 Bayes 准则推导式(7)中的子项 $p(I_i | I_{i-1}, C_i)$, 由字符图像相互独立的假设可以得到 $p(C_i | I_i, I_{i-1}) = p(C_i | I_i)$, $p(C_i | I_{i-1}) = p(C_i)$, 则有

$$p(I_i | I_{i-1}, C_i) = \frac{p(C_i | I_i, I_{i-1}) p(I_i | I_{i-1})}{p(C_i | I_{i-1})} =$$

$$\frac{p(C_i | I_i) p(I_i | I_{i-1})}{p(C_i)} = \frac{p(I_i | C_i) p(I_i | I_{i-1})}{p(I_i)} \quad (8)$$

由此式(7)可进一步写成

$$p(I_1, \dots, I_n | C_1, \dots, C_n) =$$

$$\left[\prod_{i=1}^n p(I_i | C_i) \right] \left[\prod_{i=2}^n \frac{p(I_i | I_{i-1})}{p(I_i)} \right] \quad (9)$$

式(9)右边第一项表征字符识别信息,其中 $p(I_i | C_i)$ 是字符 C_i 被识别成 I_i 的置信度。现行维文含有 128 个变体字符,本文采用对单字符提取较为流行的上、下、左、右 4 个方向的 8×8 弹性网格特征^[14],及改进的二次鉴别函数^[15]获得识别距离。识别置信度是通过分类器输出进行置信度转换获取的,转换中采用了 sigmoid 函数结合 soft-max 函数^[16]的方法。

式(9)右边第二项表征字符序列语义信息,其中

$p(I_i)$ 和 $p(I_i | I_{i-1})$ 可分别作为二阶 Markov 语言模型的字符先验概率和二元转移概率,模型参数由 2.5 万个常用维文词汇统计得到。

2.3 加权信息融合的路径寻优方法

本文算法中的切分采用了自底向上的策略,切分路径是字段匹配成为字根,字根合并得到字符,这样字符分割的问题可转化为切分路径寻优的问题。设单词过分割后得到主体字段 $\mathbf{M} = (M_1, \dots, M_N)$ 和附加字段 $\mathbf{A} = (A_1, \dots, A_M)$, 在任意一条切分路径 s 下获得的字符图像序列为 $\mathbf{C}^s = (C_1^s, \dots, C_n^s)$, $1 \leq n \leq N$, 并以 $\mathbf{I}^{s,t} = (I_1^{s,t}, \dots, I_n^{s,t})$ 表示 $\mathbf{C}^s$ 在单字符分类器下的第 t 种候选识别结果。根据最大后验概率准则,最优切分路径对应的字符序列可表述为

$$(\hat{C}_1^s, \dots, \hat{C}_n^s) =$$

$$\arg \max_s [p(I_1^{s,t}, \dots, I_n^{s,t} | M_1, \dots, M_N, A_1, \dots, A_M)] \quad (10)$$

根据概率论乘法原理,式(10)右边可写成

$$p(I_1^{s,t}, \dots, I_n^{s,t} | M_1, \dots, M_N, A_1, \dots, A_M) =$$

$$p(C_1^s, \dots, C_n^s | M_1, \dots, M_N, A_1, \dots, A_M) \cdot$$

$$p(I_1^{s,t}, \dots, I_n^{s,t} | C_1^s, \dots, C_n^s) \quad (11)$$

维文字符高、宽及纵、横比均不相等,但字符附加笔画一般固定写在主体笔画的正上或正下方。鉴于此,忽略字符形状大小,字符置信度可由主体和附加字段的匹配置信度来表述。假设各字符匹配置信度仅与字符自身相关,由此式(11)右边第 1 项可写成

$$p(C_1^s, \dots, C_n^s | M_1, \dots, M_N, A_1, \dots, A_M) =$$

$$\prod_{i=1}^n p(G(C_i^s)) \quad (12)$$

式中: $p(G(C_i^s))$ 表示字符 C_i^s 的匹配置信度。将式(9)和式(12)代入式(11),再将式(11)代入式(10)的对数表达式为

$$(\hat{C}_1^s, \dots, \hat{C}_n^s) = \arg \max_s \left\{ \sum_{i=1}^n \lg p(G(C_i^s)) + \right.$$

$$\left. \sum_{i=1}^n \lg p(I_i^{s,t} | C_i^s) + \sum_{i=2}^n [\lg p(I_i^{s,t} | I_{i-1}^{s,t}) - \lg p(I_i^{s,t})] \right\} \quad (13)$$

式(13)通过综合字符图像序列的匹配信息(取决于 $p(G(C_i^s))$)、识别信息(取决于 $p(I_i^{s,t} | C_i^s)$)和语义信息(取决于 $p(I_i^{s,t} | I_{i-1}^{s,t})$ 及 $p(I_i^{s,t})$)来挑选最优的字段匹配和字根合并路径。由于置信度估计存在偏差,所以均匀融合会降低后验概率的鉴别能力,影响分割性能的稳定性。为此,考虑各信息在路径寻优中的不同权重的作用,设立了匹配、识别和语

义信息的融合系数 $(1,\alpha,\beta)$ 。在路径代价计算中,因子 $\alpha,\beta\geq 0$ 分别决定着识别信息和语义信息相对于匹配信息的比重,当 3 种融合信息的比例达到合理要求时,分割路径寻优获得最佳结果。另外,鉴于字符由多个字根组成,利用字符 C_i^s 包含字根的个数 n_i^s 来加权路径代价,以避免最优路径倾向于较短的合并路径。本文切分路径寻优计算式如下

$$(\hat{C}_1^s,\dots,\hat{C}_n^s)=\operatornamewithlimits{argmax}_{s,t}\left\{\sum_{i=1}^n n_i^s\lg p(G(C_i^s))+\right.\\ \left.\alpha\sum_{i=1}^n n_i^s\lg p(I_i^{s,t}\mid C_i^s)+\beta\sum_{i=2}^n n_i^s[\lg p(I_i^{s,t}\mid I_{i-1}^{s,t})-\right.\\ \left.\lg p(I_i^{s,t})]\right\}\tag{14}$$

实际操作时将字符识别结果限定在前 10 个候选类别,将字根合并范围限定于连体段内,这样算法将得到简化。当然,通过路径评估分数排序,也可输出多个候选分割结果。

3 实验结果与讨论

实验中手写维文单词样本集中共有 500 类词,每类中有 30 套样本,数据来自移动终端手机。实验在 IntelE7400 CPU、2.0 GB 的微机上进行,采用留一法进行交叉验证。

3.1 算法参数对切分性能的影响

实验测试了式(14)中匹配、识别和语义信息的组合系数 $(1,\alpha,\beta)$ 对本文算法性能的影响,并以字符序列的单字符平均识别率来评价分割算法的性能,结果如图 5 所示。

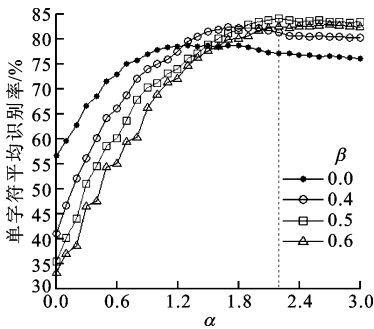


图 5 参数 α,β 对本文算法性能的影响

首先讨论 2 种特殊情况:当 $\alpha=0$ 时,字符分割路径仅由匹配信息和语义信息来决定,这时由于缺乏字符识别信息,分割准确率很低,稀有词语识别率过低对分割效果有干扰,所以 β 越大则分割效果越差;当 $\beta=0$ 时,字符分割路径取决于匹配信息和识别信息,这时算法性能较为平稳,但不考虑语义信息的作用,最优分割效果则稍差。从图 5 可以看出:对

于 α ,随着 α 的逐渐增加,识别信息对切分路径代价的影响增大,分割准确率随之上升,当 α 达到一定值时分割性能趋于稳定,当 α 过大时匹配和语义信息的比例相对减小,分割性能呈现缓慢下降的趋势;对于 β ,当 α 达到一定值时,一方面 β 增大使得语义信息增加,分割性能提升,另一方面 β 越大,分割性能趋于稳定时的 α 越大,此时匹配信息的比例过小,分割性能提升受到了限制。综合上述分析,通过实验得到经验参数 $\alpha=2.2,\beta=0.5$ 。

3.2 算法性能比较

实验对比了 4 种算法的性能。算法 1 是本文第 1 节所述的过分割算法,其中主体和附加字段以最近距离原则匹配;算法 2 是本文过分割结合多信息融合路径寻优的字符分割算法,其中 $\alpha=2.2,\beta=0.5$;算法 3 是采用算法 1 结合文献[9]的总结了切分点周围字段信息的切分路径寻优方法;算法 4 是采用算法 1 结合文献[10]的基于字符隐 Markov 模型的切分路径寻优方法。算法 2、3、4 使用了同一过分割算法,为的是便于观察路径寻优策略对分割结果的影响。实验结果如表 1 所示,其中:准确率指算法获得分割位置正确的比率;召回率指真实分割位置中被算法正确检出的比率;误检率=1-准确率,包括多余切分点和字段错误匹配 2 种情况。

表 1 算法分割性能比较

算法	准确率 /%	召回率 /%	误检率/%		每词平均耗时 /ms
			多余切分	错误匹配	
1	75.69	98.54	21.53	2.78	197
2	98.07	98.81	1.76	0.17	836
3	95.92	94.85	1.92	2.16	975
4	91.85	91.91	5.62	2.53	795

从表 1 可以看出:算法 1 在 24.31%的误检率下得到了 98.54%召回率,满足下一步切分路径寻优的需求;从整体分割效果来看,算法 2 融合了匹配、识别和语义等多种信息,并利用组合系数对各信息进行了平衡,从而获得了最高的准确率和召回率;算法 3 总结了切分点周围各字段的识别信息,分割效果较好,但算法性能仅取决于字符分类器,且神经网络分类器的运算复杂度较高;算法 4 以字符隐 Markov 模型综合了识别和语义信息,但未对这 2 种信息进行平衡,分割结果受语义信息影响较大,因而对各类单词的分割效果不一,算法鲁棒性较差。此外,算法 2 结合了字段模糊匹配和匹配路径寻优,

将过分分割算法(算法 1)中由错误匹配造成的误检率降至 0.17%,相对算法 3 和算法 4,在纠正字段误匹配上优势明显。本文算法运行结果的示例如图 6 所示,包括过分分割字根序列和最终字符分割结果。

4 结 论

本文提出了一种有效的维文字符分割算法,其中过分分割算法是整个算法的基础,它先利用基线域进行检测,结合连通域进行分析,从而有效分离了交

叠字符,再通过主体笔画切分和附加笔画进行聚类,在分割粘连字符的同时避免了附加笔画误切,进而获得了较高的过分分割召回率。切分路径寻优策略是整个算法的关键,即通过字段模糊匹配解决了笔画漂移引起的匹配错误,然后通过对匹配位置进行高斯建模来估算字符匹配信息,利用字符序列二阶 Markov 模型获取识别信息和语义信息,经匹配、识别及语义等多种信息的融合,以组合系数来调整各信息的贡献,最终获得了较为理想的分割效果。

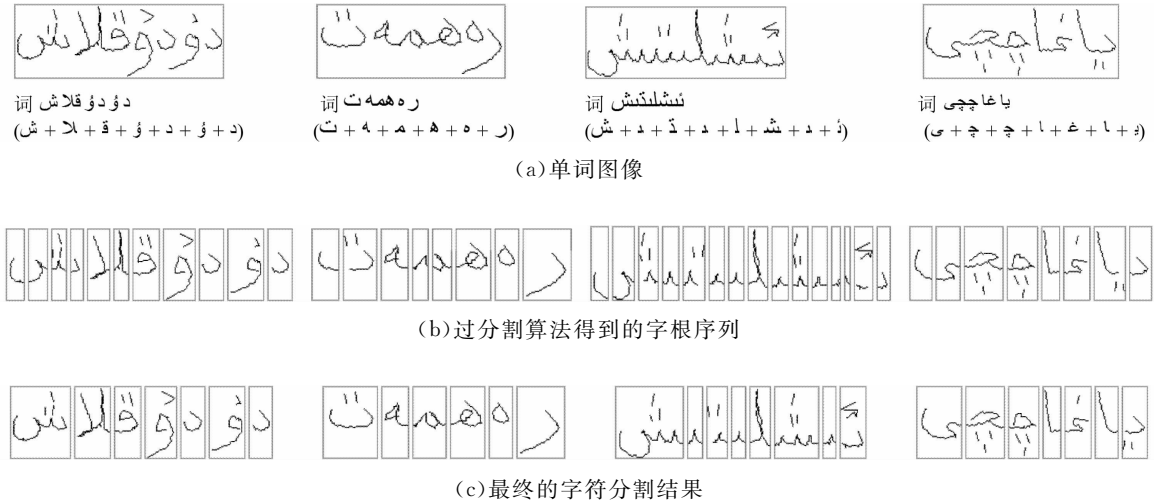


图 6 本文算法运行结果示例

参考文献:

[1] 宋喆. 现代维吾尔语词汇构成途径新探 [D]. 乌鲁木齐:新疆大学, 2006.

[2] 哈力木拉提,阿孜古丽. 多字体印刷维吾尔文字符识别系统的研究与开发 [J]. 计算机学报, 2004, 27 (11):1480-1484.

HALMURAT, AZIGULI. Research and development of a multifont printed Uyghur character recognition system [J]. Chinese Journal of Computers, 2004, 27 (11):1480-1484.

[3] 靳简明,丁晓青,彭良瑞,等. 印刷维吾尔文本切割 [J]. 中文信息学报, 2005,18(5):76-83.

JIN Jianming, DING Xiaoqing, PENG Liangrui, et al. Printed Uyghur texts segmentation [J]. Journal of Chinese Information Processing, 2005,18(5):76-83.

[4] AZEEM S A, AHMED H. Recognition of segmented online Arabic handwritten characters of the ADAB database [C]// Proceedings of the 10th IEEE International Conference on Machine Learning and Application. Piscataway, NJ, USA: IEEE, 2011:204-207.

[5] SAEED K, ALBAKOOR M. Region growing based segmentation algorithm for typewritten and handwritten text recognition [J]. Applied Soft Computing, 2009,9 (2):608-617.

[6] ABANDAH G A, JAMOUR F T. Recognizing hand-written Arabic script through efficient skeleton-based grapheme segmentation algorithm [C]// Proceedings of the 10th International Conference on Intelligent Systems Design and Applications. Piscataway, NJ, USA: IEEE, 2010:977-982.

[7] PARVEZ M T, MAHMOUD S A. Arabic handwriting recognition using structural and syntactic pattern attributes [J]. Pattern Recognition, 2013,46(1):141-154.

[8] DING Xiaoqing, LIU Hailong. Segmentation-driven off-line handwritten Chinese and Arabic script recognition [M]// Lecture Notes in Computer Science: Vol. 4768. Berlin, Germany: Springer, 2008:196-217.

[9] AL-HAMAD H A, ZITAR R A. Development of an efficient neural-based segmentation technique for Arabic handwriting recognition [J]. Pattern Recognition, 2010,43(8):2773-2798.

[10] BOUKHAROUB A, BENNIA A. Recognition of handwritten Arabic literal amounts using a hybrid approach [J]. Cognitive Computation, 2011,3(2):382-393.

(下转第 86 页)