

基于混合码本与因子分析的 文本独立笔迹鉴别

阿依夏木·力提甫^{1,2} 鄢煜尘¹ 肖进胜¹
江 昊¹ 姚渭菁³

摘 要 针对已有的笔迹鉴别方法对笔迹版式的要求比较严格、训练过程耗时、对内容不受限制的小样本数据情况下鉴别性能较低等问题,提出了基于混合码本与因子分析的文本独立笔迹鉴别算法。该算法提取写作时常用的子图像,并用描述符标注“代码”建立“码本”。在特征提取层,分别采用加权的方向指数直方图法和距离变换法,对于具有相同描述符的“代码”计算特征距离。把影响特征距离的因素分为书写因子和字符因子,对码本中的每个书写模式进行双因子方差分析。在 IAM 和 Firemaker 这两个标准数据集上的实验结果证明,相比目前国内外的先进已有方法,本文提出的算法在精度和速度方面有一定的优势,具有一定的推广价值,适合处理多语种的笔迹鉴别问题。

关键词 笔迹鉴别, 混合码本, 文本独立, 因子分析

引用格式 阿依夏木·力提甫, 鄢煜尘, 肖进胜, 江昊, 姚渭菁. 基于混合码本与因子分析的文本独立笔迹鉴别. 自动化学报, 2021, 47(9): 2276–2284

DOI 10.16383/j.aas.c190121

Text-independent Writer Identification Based on Hybrid Codebook and Factor Analysis

Ayixiamu · LITIFU^{1,2} YAN Yu-Chen¹ XIAO Jin-Sheng¹
JIANG Hao¹ YAO Wei-Qing³

Abstract In order to solve the problems of existing writer identification methods, such as strict requirements on handwriting format, time-consuming training process and low identification performance under the condition of small sample data with unlimited content, a text-independent writer identification algorithm based on hybrid codebook and factor analysis is proposed. This algorithm extracts sub-images that are often used in writing, and uses descriptors to label them as codes to create code books. In the feature extraction layer, weighted directional index histogram method and distance transformation method are used respectively to calculate feature distance for codes with the same descriptor. Then the factors that affect the feature distance are divided into writing factors and character factors, and the two-way analysis of variance is carried out for each writing mode in the codebook. The experimental results on IAM and Firemaker benchmark datasets show that compared with the current advanced methods at home and abroad, the algorithm proposed in this paper has certain advantages in accuracy and speed. Compared with some methods, it has certain promotional value and is suitable for dealing with multilingual writer identification problems.

Key words Writer identification, hybrid codebook, text independent, factor analysis

Citation Ayixiamu · Litifu, Yan Yu-Chen, Xiao Jin-Sheng, Jiang Hao, Yao Wei-Qing. Text-independent writer identification based on hybrid codebook and factor analysis. *Acta Automatica Sinica*, 2021, 47(9): 2276–2284

笔迹鉴别指的是通过手写的文字信息鉴定书写人身份的一种文件鉴定技术。它作为机器视觉与模式识别领域中近几年的研究热点之一,在历史文件分析、司法嫌疑人身份识别和古代手稿分类等方面发挥着重要作用。在过去的几十年里,笔迹专家们大都利用机器视觉技术来研究世界上主要语言的笔迹鉴别问题,然而小型语言的存在为笔迹鉴别领域提供了新的研究空间^[1]。由于每种语言脚本的独特性,各语种的笔迹鉴别技术略有不同。因每一种语言都对笔迹鉴别方法提出新的挑战,很难有适用于所有语言的通用技术。本文重点研究维吾尔文笔迹鉴别问题,并利用现有的 IAM^[2] 与 Firemaker^[3] 标准数据集验证本文算法的可行性。手写文本模式有两种:含书写文本的笔轨迹时间序列的在线模式和仅含书写文本图像的离线模式,分为在线和离线的笔迹鉴别方法^[4]。写作速度、角度、笔顺或压力用于在线笔迹鉴别,而与单词、字符、行或段落相关联的特征用于离线笔迹鉴别。本文研究的对象即为离线笔迹鉴别方法。

当前的离线笔迹鉴别方法根据提取特征方式的不同可分为全局特征提取方法^[5–6]与局部特征提取方法^[7–11]。全局特征提取方法把手写笔迹看成特殊的纹理图像,提取能够反映手写文本统计特性的全局特征作为鉴别的依据。局部特征提取方法是对笔迹图像的局部结构、梯度、轮廓、几何特征等进行特征描述,并通过编码方式将局部特征映射到公共空间形成全局特征。以往文献中提出的微结构特征^[7]局部二值模式(Local binary pattern, LBP)以及局部相位量化(Local phase quantization, LPQ)^[8]、尺度不变特征变换(Scale-invariant feature transform, SIFT)^[9–11]和高斯混合模型(Gaussian mixed model, GMM)超向量^[12]都属于局部特征提取方法。随着深度学习算法的广泛推广,基于无监督特征学习^[10]、半监督特征学习^[11]和卷积神经网络(Convolutional neural network, CNN)的笔迹鉴别^[4]方法也得到了发展。对于小样本笔迹图片,相比于全局纹理特征,笔迹的局部结构特征更直观、显著、稳定。因此,近年来大量的研究集中在基于局部结构特征的笔迹鉴别方法上,基于码本^[13–15]的笔迹特征提取是其中较重要的关注点。本文提出的方法是基于局部结构特征生成码本的方法,其主要思路是从两份笔迹文本中提取书写不变模式组成码本,然后通过提取每一个码本成员的局部特征形成全局特征。

计算机笔迹鉴别根据测试对象和特征提取的方法分为两大类:文本独立方法与文本依存方法。文本依存方法要求参考样本与测试样本的书写内容相同,并且主要依靠内

收稿日期 2019-03-01 录用日期 2019-06-02

Manuscript received March 1, 2019; accepted June 2, 2019

新疆维吾尔自治区高校科研计划自然科学基金青年项目(XJUDU2019Y032),新疆师范大学重点实验室招标课题(XJNUSYS092018A02)资助

Supported by Xinjiang Uyghur Autonomous Region University Scientific Research Program Natural Science Youth Project (XJUDU2019Y032) and Tender Subject for Key Laboratory Project of Xinjiang Normal University (XJNUSYS092018A02)

本文责任编辑 金连文

Recommended by Associate Editor JIN Lian-Wen

1. 武汉大学电子信息学院 武汉 430072 2. 新疆师范大学物理与电子工程学院 乌鲁木齐 830054 3. 国网湖北省电力有限公司信息通信公司 武汉 430077

1. Electronic Information School, Wuhan University, Wuhan 430072 2. College of Physics and Electronic Engineering, Xinjiang Normal University, Urumqi 830054 3. Information and Communication Company, State Grid Hubei Electric Power Co., Ltd., Wuhan 430077

容相同的子图像进行比较。虽然此种方法的鉴别准确率很高,但是在实践中基于固定文本的笔迹鉴别有一定的局限性。在文本独立的笔迹鉴别方法中样本的书写内容不受限制,比文本依存方法更具有广泛的应用前景。但是文本独立方法的鉴别准确性不高,并需要大量的训练样本。本文有效结合文本依存和文本独立两种方法的优点,提出了一种基于混合码本与因子分析的文本独立笔迹鉴别算法。文中首先从二值化的原始笔迹图像提取子图像并用描述符标注,引入了混合码本的概念;然后采用方向指数直方图法(Directional index histogram, DIH)和距离变换法(Distance transformation, DT)提取所有子图像的特征,计算参考样本与测试样本中具有相同描述符的子图像之间的距离。前期处理过程是典型的文本依存方法,然而本文关注的重点不在于子图像的内容,描述符只是为了快速检索相同内容的码本成员。最后通过统计学中的双因子方差分析法(Two way analysis of variance, TW-ANOVA),把影响鉴别精度的因素分为书写因子与字符因子,利用因子分离方法实现了文本独立的笔迹鉴别分类器。在分类决策层,利用特征融合与多分类器组合的方式提高笔迹鉴别准确率。在维吾尔文 2016 数据集、标准的 IAM 与 Firemaker 数据集上的实验结果表明,本文的方法只需要极少的笔迹信息就能得到较好的鉴别结果,算法运行时间短,并且相关技术可以应用于其他语种的笔迹鉴别,具有良好的应用前景和推广价值。

本文其余部分的安排如下:第 1 节为相关领域的研究现状。第 2 节详细描述了基于混合码本与因子分析的文本独立笔迹鉴别算法的流程。第 3 节给出了在维吾尔文 2016 数据集以及两个基准数据集上的实验结果与分析。第 4 节给出了结论与展望。

1 基于机器视觉的笔迹鉴别技术

如前所述,笔迹鉴别需要提取特定于书写人的笔迹特征;文本独立的笔迹特征大致可分为两类:基于纹理的全局特征和基于图形的局部特征。考虑到本文提取的笔迹特征属于局部特征提取方法,结合维吾尔文及类似文字的特点,我们将重点放在相关语言笔迹鉴别研究中表现良好的研究方法。在过去的十年中,深度学习技术成功地应用于包括笔迹鉴别在内的许多识别任务中。自从深度学习算法成功地应用于从笔迹数据中自动学习特征,以往的基于手工特征的算法被称为传统的笔迹鉴别方法。

1.1 基于传统方法的手工特征

早期文献[7]提出了从笔迹轮廓链码中提取的微结构特征用于笔迹识别,但微结构特征要从由足够篇幅的整篇文本的笔迹样本上提取,需要的样本字数相对较多,不适合实际应用。随后,纹理描述方面的算法以其快速提取纹理特征以及计算速度快等方面的优势开始普遍应用。其中,LBP 是一种灰度和旋转不变的纹理描述符,LPQ 在处理模糊纹理方面表现出很强的鲁棒性,并且在纹理分类方面优于 LBP^[8]。文献[8]提出了基于 LBP 与 LPQ 的纹理描述符提取组合纹理特征的方法,并使用相异特征向量来训练支持向量机(Support vector machine, SVM)分类器。该方法不仅解决了基于相异度的笔迹鉴别方法中存在的问题,还证明了相异度方法优于经典的分类方法。鉴于局部纹理描述符在纹理分类问题中的有效性和小书写片段在描述书写风格时的高鉴别能力,文献[16]提出了基于三种纹理描述符,即 LBP、LPQ 以及局部三元模式(Local ternary pattern, LTP)的笔迹鉴别方法。虽然文献[8]和文献[16]

获得了比较理想的笔迹鉴别效果,但是需要提取大量的书写片段,由于各种笔迹具有丰富的特征,导致书写片段之间存在局部特征相似性,从而造成的记忆限制。为解决此类问题,文献[6]提出了一种使用袋装离散余弦变换描述符的笔迹鉴别系统,离散余弦变换系数通常对书写或扫描过程中可能发生的失真具有鲁棒性。

SIFT 或类似 SIFT 的描述符是局部特征提取方法中最常见的一种,典型的 SIFT 词袋模型^[17]已经在文献[9, 18–19]中有所应用。SIFT 描述符在图像检索以及图像取证相关领域^[12]有着强大的功能,但需要组合能力强的编码方式。SIFT 类研究工作中,文献[9]通过计算不同笔迹的 SIFT 特征,使用 K 均值进行聚类搭建了词袋模型。在此基础上,文献[18]先用各向同性对数滤波器把手写图像分割成单词区域,然后提取 SIFT 特征以及相应的尺度和方向特征。文献[19]进一步提出了从图像中提取的一组 SIFT 描述符进行聚类来构建局部纹理模式的码本,然后使用轮廓方向特征和 SIFT 描述符细化候选列表的文本独立分类器。文献[12]使用在脚本轮廓处密集计算的 RootSIFT 描述符,并将 GMM 超向量用作笔迹特征的编码方法。该文使用样本 SVM 来训练特定于文档的相似性度量,扩展了文献[19]的工作。文献[20]将 SIFT 和 RootSIFT 描述符结合在一起组成了 GMM,通过加权直方图的评估,获得了很高的笔迹鉴别准确率。

最近几年来,在提取局部结构特征方面也利用基于码本的笔迹鉴别算法。文献[13]提取的码本更注重于图像的方向和曲率特征,并证明在预处理的过程中笔迹图像有任何形状变化会对鉴别准确率引起比较大的影响。文献[14]使用两种有效的轮廓码提取方法,但对于子图像的切分要求比较严格。文献[15]提出的集成码本具有多个不同大小的码本,类似于文献[14]的字符碎片码本,计算复杂度比较高。本文深入研究各语种文字的结构特征,提出了基于笔迹书写结构切分子图像的码本特征。在预处理阶段,高频模式的切分工作不受窗口大小和形状变换的影响,并且需要提取的代码数量远比以前文献少。在测试阶段使用简单易行的两种传统特征提取方法,计算量相对较少,更重要的是书写人数的增多对实验结果的影响较不明显。当书写人数增加时,本文算法有较强的鲁棒性。

总之,虽然上述文献提及的 SIFT 类描述符、离散余弦变换描述符以及其他类型的描述符都可以进行笔迹鉴别,但是比较适合用于测试样本上的字数较多的笔迹鉴别任务中。在实际应用中,经常会面临内容不受限制以及样本字数相对较少的情况。本文算法在预处理、子图像切分、特征提取等各个方面有一定的优势,具有一定的参考价值和可比性。

1.2 基于深度学习的特征

如前所述,手工特征很难做出定义,并且特征提取过程比较复杂。传统的监督学习需要大量的标号样本,而无监督的学习方法仅仅使用无标号样本。在文献[10]中提出的是以无监督的方式学习深度卷积神经网络(Deep convolutional neural network, DCNN)的激活特征方法。半监督学习即从有标号样本和无标号样本中学习。近年来,基于神经网络的技术也已应用于笔迹鉴别方面^[21–22]。这些技术利用 CNN 的优点来解决自动特征提取的问题。对笔迹鉴别任务,文献[22]采用了 CNN 作为局部特征提取器。该方法需要对图像进行二值化和归一化预处理,因此其性能取决于数据库和预处理方法。文献[21]提出了另一种策

略: 在从 CNN 提取局部特征后, 它们被用于基于 GMM 超矢量编码形成全局特征. 这种组合方法比文献 [22] 提出的方法表现得更好. 然而, 文献 [21] 和文献 [22] 有两个独立的训练步骤: 特征提取和编码, 其中 CNN 预先训练用于提取局部特征. 也就是说, 在训练和编码的第 2 步中, 预先训练的 CNN 系统是固定的, 没有更新, 降低了整个系统的性能. 因此, 文献 [4] 使用端到端的神经网络进行笔迹鉴别, 其中基于 CNN 的特征提取器和基于神经网络的分类器连接并一起训练. 虽然深度学习 (Deep learning, DL) 算法实现了自动学习笔迹特征的优势, 但其网络结构庞大, 训练权值多, 因此需要海量的训练数据进行训练, 通常需要大量带注释的训练数据. 现实应用中受存储空间、获取样本时间等限制, 往往存在训练样本不足的问题, 这将直接影响识别的准确率.

本文基于 DIH 和 DT 等经典算法, 分别提取纹理特征和结构特征, 实现过程简单易行, 不需要大量的训练样本, 对设备的要求不高, 不易受到样本数量的影响. 本文为了提高笔迹鉴别效率, 采取了纹理特征和结构特征的组合分类措施^[23], 尤其是在样本字数较少, 内容不受限制的场合更能体现本文系统的优越性, 与深度学习方法以及以往的研究方法相比, 笔迹鉴别性能有着可比性.

2 混合码本生成与因子分析

现有的大多数手写笔迹鉴别系统使用统计或基于模型的方法确认书写人身份. 本文提出一种将混合码本模型和 TW-ANOVA^[24] 相结合的方法进一步提高鉴别性能, 其流程如图 1 所示.

此流程图主要包括三个部分: 混合码本生成、特征提取和因子分析. 我们的笔迹鉴别系统分别由预处理软件和测试软件组成, 其中生成码本部分利用预处理软件实现, 特征提取、因子分离以及分类决策过程通过测试软件实现.

我们首先把所有扫描好的笔迹样本分成两大组: 参考样本和测试样本. 在混合码本生成部分, 先对所有笔迹样本进行黑白化、去除各种噪声、行线以及格线变成二值图像: 然后根据特定语言的书写特点提取高频子图像, 并归一化后用描述符标注变成代码, 建立书写人的码本. 子图像的切分是整个笔迹鉴别系统的基础, 标注是为了便于检索, 选择子图像与标注方法将在第 2.1 节描述. 在特征提取层, 先把所有的码本用于建立一个参考库, 然后利用数据挖掘技术检索具有相同描述符的代码: 对于描述符匹配的子图像分别采用加权的方向指数直方图法和距离变换法提取特征并计算特征距离, 相关内容将在第 2.2 节介绍. 在因

子分析部分, 先把影响识别精度的因素分为书写因子和字符因子, 对码本中的每个书写模式进行双因子方差分析 (TW-ANOVA), 然后滤除字符因素, 得到只保留书写因素的文本独立笔迹分类器, 经过特征融合得到书写人排序, 相关内容将在第 2.3 节介绍.

2.1 混合码本的生成

现代维吾尔语是从右向左水平书写的规范性书面语言, 维吾尔文书写系统最显著的特点是每个字母有 2~6 种书写形式, 这些字母根据单词中的位置有不同的写法, 如图 2 所示. 我们通过两个维吾尔文单词描述子图像的选择过程, 虚线框所示的为相同子图像.

每个子图像可以作为一个代码, 手写文本上的高频模式无论它是单词、字母、前缀、后缀还是中缀, 只要易于切分都可以被选取, 所以称之为混合码本. 从手写图像上切分的每一个子图像都非常重要, 我们除了注重选择具有代表意义的高频模式, 还要尽量提取冗余子图像增加相同子图像的匹配概率. 与以前类似的方法不同的是我们提取的子图像经过标注环节包含一定的语义信息, 这样才能够快速检索相同子图像. 显然, 建立码本的过程类似于文本依存的笔迹鉴别方法, 它不仅适合于维吾尔文的书写特点, 还可以推广到其他语种. 考虑到 IAM 和 Firemaker 等英文数据集上手写字符数量少以及内容不受限制等因素, 我们采用的提取代码方法类似于维吾尔文 2016 数据集. 所有的子图像将组成书写人的码本, 它是手写图像的关键因素, 因为它能够有效地代表原始数据.

在我们的系统中可以用三种方式进行子图像的切分和提取, 分别包括矩形框、曲线框和全自动分割框. 每个子图像的大小不一样, 利用细化算法将它们归一化为固定 64×64 大小的矩阵, 以确保书写工具的独立性. 经过归一化处理的子图像才会变成码本上的一个代码, 如图 3 所示. 图 3 显示了本文提出混合码本的生成过程, 包括从原始笔迹图像提取子图像、标注以及代码本的生成过程.

2.2 特征提取

文本依存的笔迹鉴别方法是依靠从参考样本与测试样本选取的几组相同子图像获得良好的识别结果. 本文从识别精度、验证错误率、稳定性和计算速度等方面比较了典型的几种方法, 选择了加权方向指数直方图法 (Weighted direction index histogram, WDIH) 和 DT^[25]. 实验表明, DIH 法的计算速度与字符的笔画点数成正比的, 是一种鉴

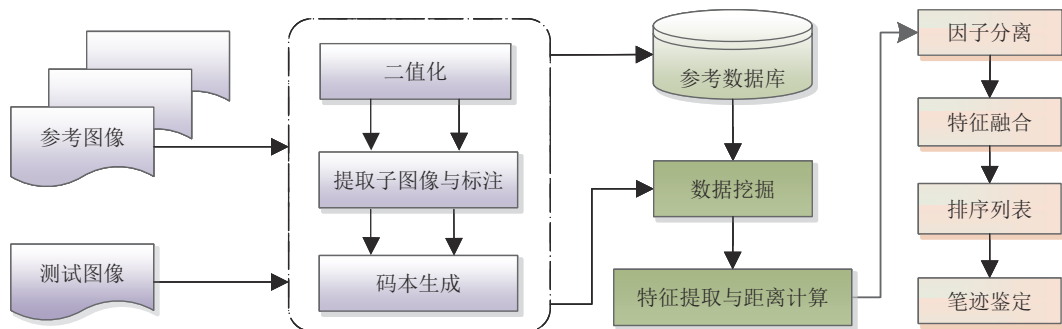


图 1 混合码本生成与因子分析的总流程图

Fig.1 The overall flow chart of proposed method

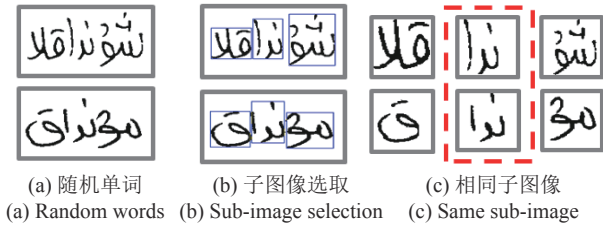


图 2 子图像的提取方法

Fig.2 Sub-image extraction method

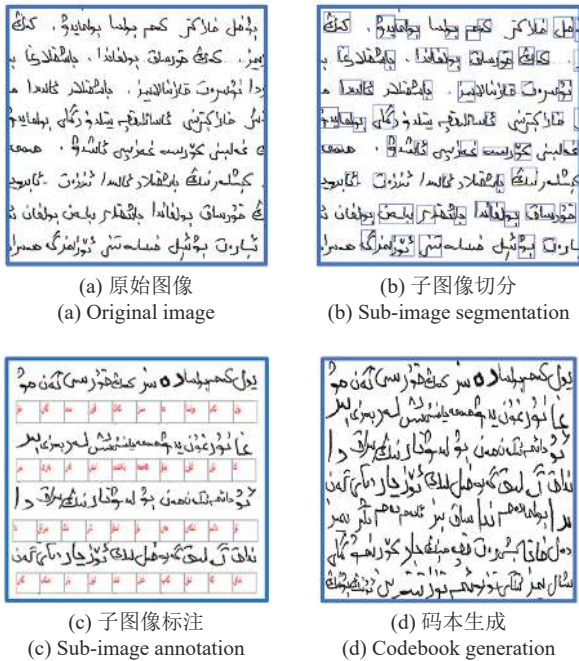


图 3 码本的生成过程

Fig.3 The generation process of codebook

别正确率高、计算速度快的鉴别方法。DT 匹配法虽然对相近模式的辨别能力不是很强,但同时也不容易把相近模式排除掉,因此实验结果表现为验证错误率较低。这两种方法的组合能够提高笔迹鉴别系统的识别率,同时能够保证系统的鲁棒性。

2.2.1 加权方向指数直方图法 (WDIH)

这是一种考虑输入图像的形状提取子图像网格特征的模板匹配方法^[26]。这种方法首先把输入图像均匀划分成 8×8 个网格,然后把每一个网格又分成 8×8 块子区域计算四个方向上的轮廓点数,得到输入图像的 8×8 个四维直方图 n_{ijk} ,其中, $i, j = 1, 2, \dots, 8$ 表示网格位置, $k = 0, 1, 2, 3$ 表示方向,获得的直方图反映了子区域中的轮廓形状。文中确定局部笔划方向的方法为:当轮廓点有一个四邻域点为零时,以该邻域点相对当前轮廓点方向的垂直方向作为笔划方向。当轮廓点有两个四邻域点为零时,若这两个邻域是连通的,以它们的联机方向作为笔划方向,否则以它们联机的垂直方向作为笔划方向。若轮廓点有三个四邻域点为零,则以不为零的那个邻域点相对当前轮廓点方向的垂直方向作为笔划方向,四个邻域点都等于零的情况则不予考虑。然后,使用均方差 $\sigma^2 = 40$ 的高斯函数对 n_{ijk} 在 8×8 的网格平面上进行空间平滑,同时采样

4×4 个点的值作为特征,链码生成 $4 \times 4 \times 4 = 64$ 位特征向量,计算式为

$$f_{uvk} = \sum_i \sum_j n_{ijk} \exp \left(-\frac{(x_i - x_u)^2 + (y_j - y_v)^2}{2\sigma^2} \right) \quad (1)$$

式中, (x_u, y_v) 表示采样点在字符图像中的坐标, (x_i, y_j) 是 8×8 网格中心点的坐标,且 $u, v = 0, 1, 2, 3$ 。得到 64 位特征矢量 f 后,计算子图像之间的距离度量 $d(f_1, f_2)$ 并进行书写人识别。

$$d(f_1, f_2) = \frac{\sum_{i=1}^{64} |f_{1i} - f_{2i}|}{\sqrt{\sum_{i=1}^{64} f_{1i} \sum_{i=1}^{64} f_{2i}}} \quad (2)$$

下面举例说明 WDIH 特征的提取过程,如图 4 所示。图 4 中输入的子图像是单词 “the”, 首先将原始图像均匀划分成 8×8 个网格,取出一个网格又分成 8×8 块子区域,并计算 4 个方向上的轮廓点数生成方向指数直方图,每个子图像总共有 $8 \times 8 \times 4$ 位方向指数直方图,采样 4×4 点后只剩下 4×4 个矩阵对应的点,并且只需要计算采样点的值。图 4(a) 中被圆圈包围的子区域根据 WDIH 的特征提取规则画出了图 4(b) 中的轮廓跟踪图,其四个方向上的方向指数直方图模型如图 4(c) 所示。

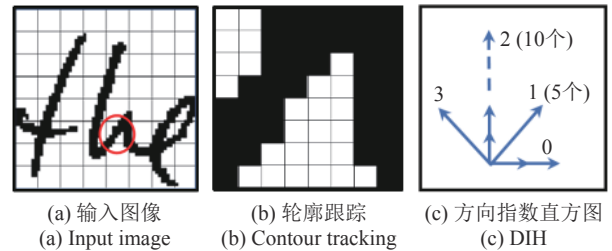


图 4 单词 “the” 的加权方向指数直方图

Fig.4 Weighted direction index histogram of “the”

2.2.2 距离变换法 (DT)

距离变换是用领域点的距离变换值来更新当前点的距离值^[25],领域是一个移动的 $k \times k$ 窗口,若领域点的值加上一个权值小于窗口中心点的值,则用这个值更新中心点的距离值。对于街区距离,当权值分别为 $a = 3, b = 4$ 时, 3×3 窗口接近欧氏距离,变换后的距离值大约是实际值的三倍。例如,图 5 所示的是数字 6 及其 DT 图,其中 3×3 窗口网格中的值是相应位置的权值。

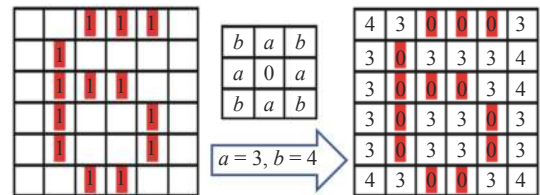


图 5 数字 “6” 的距离变换

Fig.5 Distance transformation of number “6”

设两幅图像分别表示为 $f(x, y)$, $g(x, y)$, 并且 $g(x, y)$ 的距离变换表示为 $g_d(x, y)$, 则两个图像之间的距离为

$$D_{fg} = \frac{1}{N_f} \sum_x f(x, y) g_d(x, y) \quad (3)$$

式中, N_f 是图像 f 中的黑点数量, 匹配距离与方向无关, 用同样的方法可以计算 D_{gf} .

笔迹鉴别过程中单一分类器可能存在片面性, 通过组合几种分类器可以提高分类的稳定性和准确性. 文中通过方向指数直方图法和街区距离变换法的串联组合模式进行分类, 因为方向指数直方图法的特征才 64 位, 可以在尽量不遗漏疑似笔迹的情况下, 先剔除大部分相似度较大的笔迹样本, 之后利用街区距离变换法对剩余的笔迹进行分类鉴别. 这样才能尽可能地提高鉴别速度. 分类器的组合算法比较多, 本文采用最高序号法作为多分类器组合鉴别的决策策略, 把计算出来的距离值按照与检验笔迹的相似程度从高到底的序列排序.

2.3 因子分离

本节分析因子分离的理论基础^[24], 通过实验数据和分析验证因子分离的必要性和优越性, 因子分离过程是文本依存与文本独立分类器的结合点和切换过程.

2.3.1 特征距离的影响因子分析

将影响识别精度的因素分为书写因子和字符因子两类, 其中书写因子指的是书写人书写风格的差异, 它是笔迹鉴别的基础; 而字符因子是字符形状结构的差异, 它跟笔迹内容相关, 是影响笔迹鉴别准确率的因素. 然后对特征距离进行加权变换, 去除字符因子, 得到基于文本依存方法的文本独立笔迹鉴别分类器^[25]. 文中选择几个参考样本, 分别编号为 $i = 1, 2, \dots, N$. 假设测试样本和每个参考样本具有 M 个相同的单词, 编号为 $j = 1, 2, \dots, M$. 参考样本 i 和测试样本中编号为 j 的子图像之间的距离是 d_{ij} , 其概率分布近似正态分布, 即有 $d_{ij} \sim N(\mu_{ij}, \sigma^2)$. 若假设存在总变量 μ , 编号为 i 的笔迹的书写因子为 μ_i , 且编号为 j 的子图像的字符因子为 μ_j , 则有

$$\begin{cases} \mu = \frac{1}{MN} \sum_{i=1}^N \sum_{j=1}^M \mu_{ij} \\ \mu_i = \frac{1}{M} \sum_{j=1}^M \mu_{ij}, \quad i = 1, 2, \dots, N \\ \mu_j = \frac{1}{N} \sum_{i=1}^N \mu_{ij}, \quad j = 1, 2, \dots, M \end{cases} \quad (4)$$

此时, 可以建立影响特征距离的书写因素和字符因素的方差分析模型 d_{ij} , 而不考虑书写因素和字符因素之间的交互作用. 此模型可以用下式表示为

$$d_{ij} = \mu + \alpha_i + \beta_j + \varepsilon_{ij} \quad (5)$$

其中, $\alpha_i = \mu_i - \mu$ 是参考样本中编号为 i 的笔迹的书写因子对 d_{ij} 的影响, 且 $\sum_i \alpha_i = 0$; $\beta_j = \mu_j - \mu$, 这里 β_j 是编号为 j 的子图像对 d_{ij} 的影响. 有 $\sum_j \beta_j = 0$; ε_{ij} 是平均值为 0 的随机误差, 有 $\varepsilon_{ij} \sim N(0, \sigma^2)$.

若引入记号 $\mu = \bar{d}$, 则有总平方和 S_T , 书写因子效应平方和 S_A , 字符因子效应平方和 S_B , 则 S_T , S_A 及 S_B 的定义为

$$\begin{cases} S_T = \sum_{i=1}^N \sum_{j=1}^M (d_{ij} - \bar{d})^2 \\ S_A = M \sum_{i=1}^N (\bar{d}_i - \bar{d})^2 \\ S_B = N \sum_{j=1}^M (\bar{d}_j - \bar{d})^2 \end{cases} \quad (6)$$

误差平方和定义为

$$S_E = S_T - S_A - S_B \quad (7)$$

显然, S_A 与 S_B 服从 χ^2 分布, 两个 χ^2 分布的比值为 F 分布. 可以定义如下两个 F 分布

$$\begin{cases} F_A = \frac{\frac{S_A}{N-1}}{\frac{S_E}{(N-1)(M-1)}} \\ F_B = \frac{\frac{S_B}{M-1}}{\frac{S_E}{(N-1)(M-1)}} \end{cases} \quad (8)$$

若 $F(M, N)$ 表示显著水平为 α 、自由度为 (M, N) 的 F 分布, 可以得到书写因素与字符因素的显著性假设成立的条件为

$$\begin{cases} F_A \geq F_{\alpha}(N-1, (N-1)(M-1)) \\ F_B \geq F_{\alpha}(M-1, (N-1)(M-1)) \end{cases} \quad (9)$$

2.3.2 特征距离的双因子显著性假设实验

选取由 11 ($N = 11$) 人书写的 21 个不同的单词, 共计 210 ($M = 210$) 个字符, 去除诸如斑点和网格线的噪声之后, 获得归一化的字符图像, 如图 6 所示. 图 6 左列为一列机打单词, 书写人根据行头的单词抄写 10 次即可; 图 6 上方一行数字表示单词的编号.



图 6 方差分析笔迹图像

Fig. 6 Handwriting image of variance analysis

这里可以通过提取所有子图像的方向指数特征和距离变换特征并计算特征距离来获得方差分析结果, 表 1 显示了双因子方差分析实验需要的变量和公式, 表 2 显示了两种方法的实验结果.

表 1 双因子方差分析 (TW-ANOVA) 指示表

Table 1 Two way analysis of variance instruction table

方差来源	平方和	自由度	均方	F 比
书写因子	S_A	$N - 1$	$S_A / (N - 1)$	F_A
字符因子	S_B	$M - 1$	$S_B / (M - 1)$	F_B
误差	S_E	$(N - 1)(M - 1)$	$S_E / ((N - 1)(M - 1))$	
总和	S_T	$MN - 1$		

表 2 加权方向指数直方图法/距离变换法的
TW-ANOVA 结果

方差来源	平方和	自由度	均方	F比
书写因子	1.76/4.23	10	0.176/0.423	24.11/34.67
字符因子	28.14/31.52	209	0.1346/0.1495	18.44/12.25
误差	15.23/25.61	2 090	0.0073/0.0122	
总和	45.13/61.36	2 309		

对于自由度分别为 (10, 2090) 和 (209, 2090) 的 F 分布可以观察不同 α 水平上的值, α 与 $F_{\alpha}(10, 2090)$ 以及 α 与 $F_{\alpha}(209, 2090)$ 之间的关系如图 7 所示。

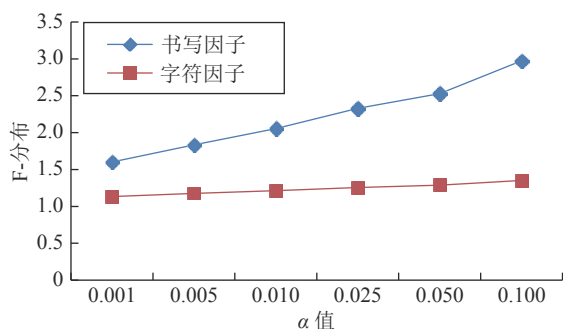


图 7 α 与 $F_{\alpha}(10, 2090)$ 和 $F_{\alpha}(209, 2090)$ 之间的关系

Fig.7 The relationship between α and $F_{\alpha}(10, 2090)$ and $F_{\alpha}(209, 2090)$

从图 7 可见, 无论 α 如何取值, $F_{\alpha}(10, 2090)$ 和 $F_{\alpha}(209, 2090)$ 始终小于 3, 表明对于书写因子 F_A 与字符因子 F_B 满足书写因子与字符因子的显著性假设成立的条件. 可以验证, 字符因子和书写因子对特征距离的影响是非常显著的。

实验结果表明, 我们可以利用因子分离法去除字符因子的影响, 才能得到与文本无关的笔迹分类器. 理想情况下有 $\mu_j = \bar{d}_j$, 因此有分类式 (10):

$$\tilde{d}_{ij} = d_{ij} - \bar{d}_j = \alpha i + \varepsilon_{ij} \quad (10)$$

可见, 距离度量 $\tilde{d}_{ij}$ 中仅含有书写因子, 因而是文本独立的, 式 (11) 中的 $\tilde{d}_i$ 可以作为文本独立的笔迹分类器。

$$\tilde{d}_i = \sum_{j=1}^M (d_{ij} - \bar{d}_j), \quad i = 1, 2, \dots, N \quad (11)$$

由式 (11) 可知, 因子分离后的子图像跟其内容无关, 测试样本与每一份参考样本中能够匹配的子图像不一样, 特征距离是经过集成化处理后的综合距离。

3 实验结果

为了验证本文算法, 并与之前的研究工作进行比较, 文中使用了维吾尔文 2016 数据集、英文 Firemaker 和 IAM 手写文本数据集。

1) 维吾尔文 2016 数据集. 此数据集是由本文作者收集的维吾尔文数据集. 为了收集符合研究要求的维吾尔文笔迹样本, 作者组织 180 名年龄在 15~70 岁之间的维吾尔族人, 并按照指定的 20 个题目, 在 A4 纸上随意书写字数不少于 50 个单词的两页文字, 每一份样本分别以 300 dpi

的分辨率扫描, 分配唯一的文件名, 并以 256 灰度级及 BMP 格式存储文件建立数据集, 后来此数据集命名为维吾尔文 2016 数据集. 该数据集的书写人性别、年龄比例相等, 包括各种教育背景的人, 书写内容相对全面、接近于真实场景, 基本满足论文需求. 测试过程中, 把同一作者提供的两页文字分成两组, 分别用于训练和测试。

2) IAM 数据集. IAM 数据集是在手写识别和书写者识别等问题上最著名和广泛使用的英文数据集之一. 它包括一些 300 dpi、8 位/像素灰度、内容各异的手写英文文本, 此数据集共包括 657 名作者的手稿, 其中 356 名作者只有一页, 301 名作者至少有两页, 125 名作者至少有四页. 对于包括两页及以上文字的样本只保留前两页, 第 1 页用于训练/验证, 第 2 页用于测试. 对于只提供一页字的作者来说, 所提供的页面大致分为两半: 前半部分用于训练/验证, 后半部分用于测试. 因此, 356 名作者有半页纸, 其他 301 名作者有一页纸用于训练/验证。

3) Firemaker 数据集. 对于 250 名书写人提供的 Firemaker 数据集, 包括根据不同的需求收集的四个子集. 本文只使用其中的第 1 个子集, 该子集包含使用普通手写的文本复制页面, 每位书写人只提供了一份样本. 同样在我们的实验中该页面被分为两部分, 分别用作参考样本和测试样本。

3.1 代码数量与书写人数对笔迹鉴别准确率的影响

这部分通过设计两种实验, 分别测试代码数量与书写人数的变化对笔迹鉴别准确率的影响. 测试目的是确认本文算法对于代码数量的最低要求以及书写人数的增多对鉴别精度的影响。

3.1.1 代码数量对鉴别精度的影响实验

我们对来自 IAM 数据集中贡献了至少两页的 180 名作者进行了代码数量对笔迹鉴别准确率的影响实验. 从每一份样本提取的代码数量从 3 增加到 70 时, 参考样本与测试样本之间的相同子图像数量从 0 增加到 33, 实验结果如图 8 所示。

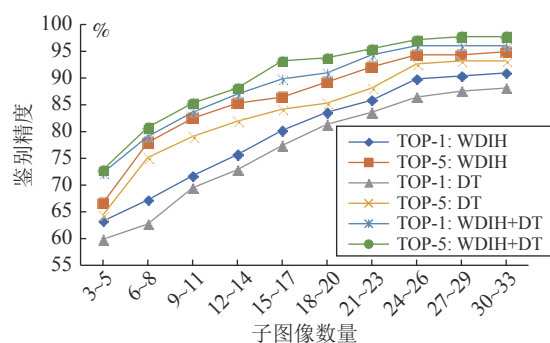


图 8 子图像数量与鉴别准确率之间的关系

Fig.8 Relationship between number of codes and identification accuracy

图 8 中以 WDIH 代表加权方向指数直方图法, DT 代表距离变换法, 这里 TOP-1, TOP-5 分别代表 1 候选和 5 候选书写人. 从图 8 可见, 凭借从参考样本提出来的 3~5 个子图像仍然可以确定书写人的身份, 但是当子图像的数量大约达到 25 幅时, 基于子图像的书写人识别率相对稳定. 此外, 与两种特征提取算法相比, WDIH 方法的性能对于子图像的数量更加敏感: 比起单一的算法两种方法的结合

可以有效提高笔迹鉴别准确率。

3.1.2 书写人数对鉴别精度的影响实验

假设从每份样本大约提取 50 幅子图像, 并且把 IAM 数据集上的书写人数量从 10 逐渐增加到 650 人时, 可以获得如图 9 所示的实验结果。

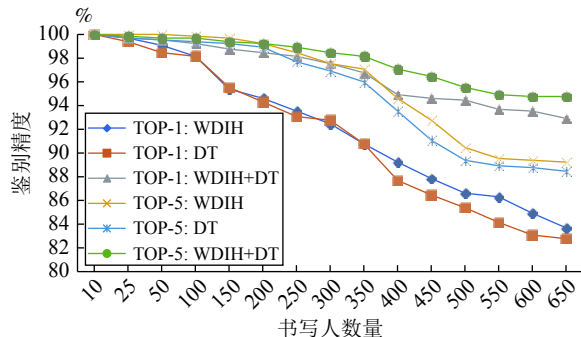


图 9 书写人数与鉴别准确率之间的关系

Fig.9 Identification accuracy with different number of writers

笔迹识别率随着书写人数量的增加持续下降, 当人数从 10 人增加到 650 人时, 三种方法的 TOP-1 识别率从 100% 分别降到 82%, 83% 和 93%。在 650 名书写人的条件下, TOP-5 的表现比 TOP-1 稳定很多, 分别下降到 88%, 89% 和 95%, 相比于 TOP-1 的鉴别率分别高于 6%, 6% 和 2%。对于两个分类器的组合模式, 虽然同样随着书写人数量的增加而出现了下降的趋势, 但是与单分类器相比, 其鉴别性能显著高于单个分类器, 并且保持相对稳定的值。

以上实验结果表明, 笔迹鉴别精度很大程度上由子图像数量与书写人数等两个因素决定。图 8 显示书写人数固定为 180 人时通过逐渐增多代码数量的方法提高了鉴别准确率。当从每一份样本提取的子图像数量大概为 30~70 个时, 能够保证系统的鲁棒性, 有效降低人数对于鉴别精度的影响。同样, 当书写人数为 650 人时, 若子图像数量从 50 幅增加到 70~100 幅, 系统的鉴别准确率则有很大的提升空间。

3.2 维吾尔文 2016 数据集上的实验结果

根据第 3.1 节的参数, 本节将展示维吾尔文 2016 数据集上的实验结果。将维吾尔文数据集中的 180 份样本人为地分成两份笔迹, 一份作为参考样本, 另一份是用于测试样本。如上所述, 为了确保测试样本与参考样本之间一定数量的相同子图像, 本文通过码本中的冗余模式提高代码之间的匹配率。因此, 我们先从每个样本中随机提取 30~40 个子图像生成码本进行测试, 然后根据每份样本的性能逐渐增加提取的子图像数量。可见, 每一份样本的码本包括很多冗余信息, 只有参考码本中的一部分代码跟测试码本上的某些代码匹配参与测试。维吾尔文 2016 数据集上的实验结果如图 10 所示, 由图可见代码数量与笔迹识别率之间的关系, 当从参考样本提取的混合子图像数量达到 70 幅时, 系统内部代码的实际匹配对会达到 20~25 个, 笔迹鉴别准确率达到理想值并且保持相对稳定。当子图像数量达到 100 个时, 系统的 TOP-1 鉴别准确率会达到 100%, 特征提取和测试过程大概需要 10~15 s。

3.3 IAM 和 Firemaker 数据集上的实验结果

本小节对现有的一些笔迹鉴别技术进行比较。本文使

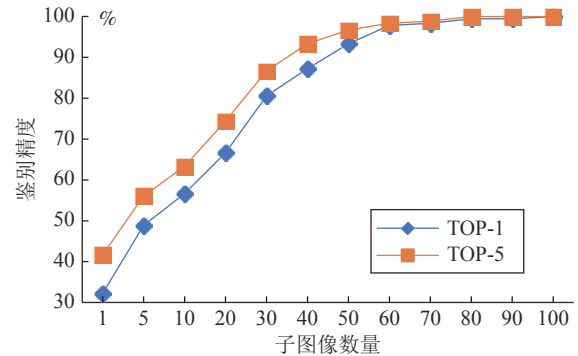


图 10 维吾尔文 2016 数据集的性能示意图

Fig.10 Performance on Uyghur2016 dataset

用应用最为广泛的两种公开数据集 IAM 和 Firemaker 作为测试数据集。为了评估模型的鲁棒性与泛化能力, 将广泛应用于笔迹检索任务中的评估标准有平均准确率均值 (Mean average precision, mAP), Soft TOP- k (TOP- k), Hard TOP- k 等几种方法^[11]。此外, 测试方法也有比较典型的几种对比策略: 一对一对比、成对对比^[14]以及相异特征对比^[8,16]等, 其中一对一对比法比较广泛应用。为了与其他文献保持一致, 本文测试过程采用一对一对比策略, 使用 TOP- k 标准用于鉴别任务中。测试过程仍然先从每个样本中随机抽取约 20 幅子图像进行测试, 然后对于鉴别失败的样本增加子图像, 并重新测试求平均值。在评估过程中, 虽然本文方法对书写页面的大小和书写字符数量的多少没有过高的要求, 但是从单个样本中提取的子图像数量要从 20 幅逐渐增加到 45 幅反复测试。具体测试结果见表 3 和表 4。表 3 的结果表明, 本文算法在 Firemaker 数据集上的 TOP-1 效果最好, TOP-10 效果与文献 [18] 相同, 均排第一。

表 3 各种方法在 Firemaker 数据集上的性能对比 (%)

Table 3 Performance comparison on Firemaker (%)

评估标准	TOP-1	TOP-10
Ghiasi (2013) ^[14]	89.2	98.6
Wu (2014) ^[18]	92.4	98.8
He (2017) ^[6]	86.2	96.6
Nguyen (2019) ^[4]	92.38	97.67
本文方法	94.4	98.8

表 4 显示本文方法在 IAM 数据集上的 TOP-10 性能最好, TOP-1 性能仅次于文献 [6] 和文献 [18] 的结果, 整体上来说效果较好。其中文献 [13~16] 采取的建立码本方法以及纹理描述符类似于本文的混合码本及其描述符, 对于 IAM 数据集, 本文结果比同类研究成果高于 1.9%~6.15%。文献 [18] 把英文 MIMUnipen 数据集用作训练数据集, 使用 IAM 和 Firemaker 作为测试数据集, 而本文和其他文献的训练和测试数据集是同一个数据集。文献 [6] 侧重研究的是系统的鲁棒性问题, 笔迹扫描质量最佳的情况下可以获得 97.2% 的 TOP-1 鉴别准确率, 但笔迹有噪声或者歪曲的条件下, 精度会下降到 92.3%, 同时样本数量增多时需要重新建立基于谱回归的核判别分析预测模型。

文献 [27] 采取一种动态片段加权组合规则减少不一致

表 4 各种方法在 IAM 数据集上的性能对比 (%)
Table 4 Performance comparison on IAM dataset (%)

评估标准	TOP-1	TOP-10
Siddiqi (2010) ^[13]	91.0	97.0
Ghiasi (2013) ^[14]	93.7	97.7
Bertolini (2013) ^[8]	88.3	—
Wu (2014) ^[18]	98.5	99.5
Khalifa (2015) ^[15]	92.0	—
Hannad (2016) ^[16]	89.54	96.77
Khan (2017) ^[6]	97.2	—
He (2017) ^[5]	89.9	96.9
Nguyen (2019) ^[4]	90.12	97.82
Hadjadji (2018) ^[27]	94.51	—
Chahi (2019) ^[28]	88.3	—
本文方法	95.69	99.69

测试片段的影响, TOP-1 笔迹鉴别率比本文结果低 1.18%.

文献 [28] 采用局部纹理特征 LBP、LTP 和 LPQ 的最佳组合模式, 得到 88.3% 的 TOP-1 笔迹鉴别率, 比本文结果低 9.72%.

本文提出的方法在预处理阶段不受窗口大小的影响, 需要切分的子图像数量相对其他方法较少, 并且书写人数的增多对实验结果的影响相对较不明显, 从书写人数增加的鲁棒性来说, 本文算法有一定优势. 本文还可以通过增多子图像数量进一步提高鉴别精度.

为了对实验方法和实验结果进行更进一步对比, 有必要讨论本文算法在不同数据集上的性能. 由表 3、表 4 及图 10 的实验结果可得本文在维吾尔文 2016, IAM 以及 Firemaker 三个数据集上的测试结果如表 5 所示.

表 5 在三个数据集上的性能对比 (%)
Table 5 Performance comparisons on three datasets (%)

评估标准	TOP-1	TOP-10
维吾尔文 2016 数据集	100	100
IAM 数据集	95.69	99.69
Firemaker 数据集	94.4	98.8

从表 5 可见基于维吾尔文的书写人识别结果高于 IAM 和 Firemaker 数据集. 出现此结果的主要原因可以归结为以下两点: 首先, 本文作者收集的维吾尔文数据集内容丰富, 字数充足. 维吾尔文 2016 数据集上的每一位书写者提供了两页维吾尔文字, 预处理阶段不仅能够快速提取高频子图像, 而且能够提取足够多的高频成分. 测试阶段, 测试样本与参考样本之间标注相同的代码数量远比 IAM 和 Firemaker 数据集的高. 因为以上英文数据集总共包括 907 个人的书写样本, 有些样本上的字数只有 4~5 行, 没有足够多的字符. 虽然预处理过程中勉强提取 30~50 个子图像, 但是实际匹配率很低. 另外, Firemaker 数据集上的手写稿包含的是固定内容, 有一定的片面性. 预处理阶段, 为了尽量选取足够多的子图像, 大部分样本按字母或者字母碎片提取代码. 虽然这种选择方法能够提高实际匹配率, 但是每一个子图像携带的笔迹特征极少, 一定程度上影响鉴别准确率. 其次, 本文提出的方法更适合于维吾尔文字

的书写特点和语法结构. 维吾尔文中字母、单词和音节的重复频率比较高. 对于本文算法, 少数子图像足以确定书写人的身份.

4 结束语

本文提出了一种用于笔迹鉴别的混合码本模型, 为了提高相同代码之间的匹配率, 此码本包括很多冗余的子图像. 对于已生成的码本先利用因子分析法, 滤除与子图像内容相关的字符因素, 保留了书写因子. 然后利用加权指数直方图法和距离变换法提取特征, 在分类决策层采用了两种方法的组合模型提高了笔迹鉴别准确率. 此外, 本文利用荷兰文和英文数据集对该方法进行了评估, 并深入研究了码本大小和书写人数对实验结果的影响. 各类实验结果表明, 我们提出的算法对于内容不受限制且字数较少的样本是非常有效的, 并且通过增加码本中的子图像数量, 可以进一步提高笔迹鉴别效率. 与 IAM 和 Firemaker 数据集相比, 在维吾尔文 2016 数据集上的实验结果非常理想, 这一结果的主要原因是维吾尔文 2016 数据集上的样本内容丰富, 并且本文算法充分利用维吾尔语的优势, 生成的码本上有足够多的子图像. 子图像数量和长度是决定本文算法鉴别准确率的关键因素.

References

- 1 Tan G J, Sulong G, Rahim M S M. Writer identification: A comparative study across three world major languages. *Forensic Science International*, 2017, **279**: 41–52
- 2 Marti U V, Bunke H. The IAM-database: An English sentence database for off-line handwriting recognition. *International Journal on Document Analysis and Recognition*, 2002, **5**: 39–46
- 3 Bulacu M, Schomaker L, Vuurpijl L. Writer identification using edge-based directional features. In: Proceedings of the 7th International Conference on Document Analysis and Recognition (ICDAR'03). Edinburgh, UK: IEEE, 2003. 937–941
- 4 Nguyen H T, Nguyen C T, Ino T, Indurkha B, Nakagawa M. Text-independent writer identification using convolutional neural network. *Pattern Recognition Letters*, 2019, **121**: 104–112
- 5 He S, Schomaker L. Writer identification using curvature-free features. *Pattern Recognition*, 2017, **63**: 451–464
- 6 Khan F A, Tahir M A, Khelifi F, Bouridane A, Almotaeryi R. Robust off-line text independent writer identification using bagged discrete cosine transform features. *Expert Systems with Applications*, 2017, **71**: 404–415
- 7 Li Xin, Ding Xiao-Qing, Peng Liang-Rui. Writer identification based on improved microstructure features. *Acta Automatica Sinica*, 2009, **35**(9): 1199–1208
(李昕, 丁晓青, 彭良瑞. 一种基于微结构特征的多文种文本无关笔迹鉴别方法. 自动化学报, 2009, **35**(9): 1199–1208)
- 8 Bertolini D, Oliveira L S, Justino E, Sabourin R. Texture-based descriptors for writer identification and verification. *Expert Systems with Applications*, 2013, **40**(6): 2069–2080
- 9 Fiel S, Sablatnig R. Writer retrieval and writer identification using local features. In: Proceedings of the 10th IAPR International Workshop on Document Analysis Systems. Gold Coast, QLD, Australia: IEEE, 2012. 145–149
- 10 Christlein V, Gropp M, Fiel S, Maier A. Unsupervised feature learning for writer identification and writer retrieval. In: Proceedings of the 14th IAPR International Conference on Docu-

- ment Analysis and Recognition (ICDAR' 17). Kyoto, Japan: IEEE, 2017. 991–997
- 11 Chen Shi-Ming, Wang Yi-Song. A robust off-line writer identification method. *Acta Automatica Sinica*, 2020, **46**(1): 108–116 (陈使明, 王以松. 一种鲁棒的离线笔迹鉴别方法. *自动化学报*, 2020, **46**(1): 108–116)
 - 12 Christlein V, Bernecker D, Hönig F, Maier A, Angelopoulou E. Writer identification using GMM supervectors and exemplar-SVMs. *Pattern Recognition*, 2017, **63**: 258–267
 - 13 Siddiqi I, Vincent N. Text-independent writer recognition using redundant writing patterns with contour-based orientation and curvature features. *Pattern Recognition*, 2010, **43**(11): 3853–3865
 - 14 Ghiasi G, Safabakhsh R. Offline text-independent writer identification using codebook and efficient code extraction methods. *Image and Vision Computing*, 2013, **31**(5): 379–391
 - 15 Khalifa E, Al-Maadeed S, Tahir M A, Bouridane A, Jamshed A. Off-line writer identification using an ensemble of grapheme codebook features. *Pattern Recognition Letters*, 2015, **59**: 18–25
 - 16 Hannad Y, Siddiqi I, El Kettani M E Y. Writer identification using texture descriptors of handwritten fragments. *Expert Systems with Applications*, 2016, **47**: 14–22
 - 17 Lowe D G. Distinctive image features from scale-invariant key points. *International Journal of Computer Vision*, 2004, **60**(2): 91–110
 - 18 Wu X Q, Tang Y B, Bu W. Offline text-independent writer identification based on scale invariant feature transformation. *IEEE Transactions on Information Forensics and Security*, 2014, **9**(3): 526–536
 - 19 Xiong Y J, Wen Y, Wang P S P, Lu Y. Text-independent writer identification using SIFT descriptor and contour-directional feature. In: Proceedings of the 13th International Conference on Document Analysis and Recognition (ICDAR' 15). Tunis, Tunisia: IEEE, 2015. 91–95
 - 20 Khan F A, Khelifi F, Tahir M A, Bouridane A. Dissimilarity Gaussian mixture models for efficient offline handwritten text-independent identification using SIFT and RootSIFT descriptors. *IEEE Transactions on Information Forensics and Security*, 2019, **14**(2): 289–303
 - 21 Christlein V, Bernecker D, Maier A, Angelopoulou E. Offline writer identification using convolutional neural network activation features. In: Proceedings of the 2015 German Conference on Pattern Recognition. Cham: Springer, 2015. 540–552
 - 22 Fiel S, Sablatnig R. Writer identification and retrieval using a convolutional neural network. In: Proceedings of the 2015 International Conference on Computer Analysis of Images and Patterns. Cham: Springer, 2015. 26–37
 - 23 He S, Schomaker L. Deep adaptive learning for writer identification based on single handwritten word images. *Pattern Recognition*, 2019, **4**(88): 64–74
 - 24 DeGroot M H, Schervish M J. *Probability and Statistics*. Beijing: Higher Education Press, 2005. 324–332
 - 25 Yan Yu-Chen, Li Cai-Yuan, Qiu Yi-Ming, Chen Qing-Hu. Text-independent classifier for handwriting verification based on factor analysis. *Engineering Journal of Wuhan University*, 2018, **51**(1): 91–94 (鄢煜尘, 李蔡媛, 邱益鸣, 陈庆虎. 基于因子分析的文本独立笔迹鉴定分类器. *武汉大学学报(工学版)*, 2018, **51**(1): 91–94)
 - 26 Xiao J S, Tian H, Zhang Y Q, Zhou Y Q, Lei J F. Blind video denoising via texture-aware noise estimation. *Computer Vision and Image Understanding*, 2018, **169**: 1–13
 - 27 Hadjadji B, Chibani Y. Two combination stages of clustered one-class classifiers for writer identification from text fragments. *Pattern Recognition*, 2018, **82**: 147–162
 - 28 Chahi A, Merabet Y EI, Ruichek Y, Touahni R. An effective and conceptually simple feature representation for off-line text-independent writer identification. *Expert Systems with Applications*, 2019, **123**: 357–376
- 阿依夏木·力提甫** 武汉大学电子信息学院博士研究生. 2012 年获得南京理工大学光电学院工学硕士学位. 主要研究方向为图像处理与模式识别. E-mail: Ayixia@whu.edu.cn
- (Ayixiamu · LITIFU)** Ph.D. candidate at the Electronic Information School, Wuhan University. She received her master degree from Nanjing University of Technology in 2012. Her research interest covers image processing and pattern recognition.)
- 鄢煜尘** 2009 年于武汉大学获工学博士学位. 主要研究方向为图像处理与模式识别. E-mail: yyc@whu.edu.cn
- (YAN Yu-Chen)** Received his Ph.D. degree in engineering from the Electronic Information School, Wuhan University in 2009. His research interest covers image processing and pattern recognition.)
- 肖进胜** 武汉大学电子信息学院副教授. 2001 年于武汉大学获理学博士学位. 主要研究方向为视频图像处理, 计算机视觉. E-mail: xiaojsh@whu.edu.cn
- (XIAO Jin-Sheng)** Associate professor at the Electronic Information School, Wuhan University. His research interest covers video and image processing, and computer vision.)
- 江 昊** 博士, 武汉大学电子信息学院教授. 主要研究方向为移动自组网络, 移动大数据, 数据挖掘. 本文通信作者. E-mail: jh@whu.edu.cn
- (JIANG Hao)** Ph.D., professor at the Electronic Information School, Wuhan University. His research interest covers mobile ad hoc network, mobile big data, and data mining. Corresponding author of this paper.)
- 姚渭箐** 国网湖北省电力有限公司信息通信公司工程师. 2017 年于武汉大学获工学博士学位. 主要研究方向为网络通信, 视频图像处理. E-mail: ywq1005@whu.edu.cn
- (YAO Wei-Qing)** Engineer at the Information Telecommunication Company State Grid Hubei Electric Power Co., Ltd. She received her Ph.D. degree in engineering from the Electronic Information School, Wuhan University in 2017. Her research interest covers network communication, video and image processing.)