

基于半监督深度信念网络的图像分类算法研究

朱常宝 程 勇 高 强
(北京化工大学信息学院 北京 100029)

摘 要 近年来,深度学习在图像、语音、视频等非结构化数据中获得了成功的应用,已成为机器学习和数据挖掘领域的研究热点。作为一种监督学习模型,成功的深度学习应用往往要求较大的高质量的训练集。基于此,研究了多个受限波尔兹曼机组成的深度信念网络,结合半监督学习的思想,使用较小的训练集提高深度网络模型分类准确性。分别采用了 Knn, SVM 和 pHash 3 种方法来学习非标示数据集,实验结果表明半监督深度信念网络比传统多层受限波尔兹曼机在图像分类准确率方面提高了约 3%。

关键词 半监督学习,深度信念网络,受限波尔兹曼机

中图法分类号 TP391.1 文献标识码 A

Research on Image Classification Algorithm Based on Semi-supervised Deep Belief Network

ZHU Chang-bao CHENG Yong GAO Qiang

(Information Institute, Beijing University of Chemical Technology, Beijing 100029, China)

Abstract In recent years, the deep learning to get a successful application in image, voice, video and other unstructured data, has become a hot topic of machine learning and data mining. As a supervised learning model, the successful deep learning applications often require a larger set of high-quality training. Based on this situation, we studied deep belief network composed of more restricted Boltzmann machines, and combined with the thought of semi-supervised learning, we used smaller training set to improve the classification accuracy of depth network model. We used Knn, SVM and pHash three methods to study the non-labeled data set. And the result shows that the semi-supervised deep belief networks increases image classification accuracy by about 3% compared with the traditional network with more restricted Boltzmann machine.

Keywords Semi-supervised, Deep belief networks, Restricted boltzmann machine

1 引言

近年来,深度学习(Deep Learning)在语音识别^[1,2]、图像分类、中文文本分类、情感识别等非结构化数据中获得了巨大的成功,受到机器学习和数据挖掘等领域的广泛重视,成为当前的研究热点。现在国内知名互联网公司都在大力地发展深度学习,例如 Google 公司的谷歌大脑(Google Brain)项目,通过将 1.6 万台电脑的处理器相连接建造出了全球为数不多的最大中枢网络系统,并具有自主学习能力。国内百度公司也大力开发了百度大脑项目,聘请了国外知名教授吴恩达来推动深度学习的发展。2015 年 3 月两会期间,百度 CEO 李彦宏还建议设立“中国大脑”计划,推动人工智能跨越发展,抢占新一轮科技革命制高点。

深度学习的概念最初由加拿大多伦多大学 Geoffrey Hinton 等人首次提出^[3,4],其基本思想来源于人工神经网络,例如含多个隐含层的多层感知器就是一种深度学习结构,其他常用模型或方法包括自动编码器、稀疏编码、深度置信网络

和卷积神经网络等。相比传统的浅层学习,这种结构的一个突出优点是解决了如何处理“抽象概念”这个亘古难题。当前多数分类、回归等学习方法为浅层结构算法,其局限性在于在有限样本和计算单元的情况下对复杂函数的表示能力有限,针对复杂分类问题,其泛化能力受到一定的制约。深度学习可通过学习一种深层非线性网络结构实现复杂函数逼近,表征输入数据分布式表示,并展现了强大的从少数样本中学习数据集本质特征的能力。深度学习的实质是通过构建具有很多隐层的神经网络学习模型和海量的训练数据来学习更有用的特征,从而最终提升分类或预测的准确性。因此“深度模型”是手段,“特征学习”是目的。区别于传统的浅层学习,深度学习的不同点在于:1)强调了模型结构的深度,通常有 5 层、6 层,甚至 10 多层的隐层节点;2)明确突出了特征学习的重要性,即通过逐层特征变换,将样本在原来空间的特征表示转换到一个新特征空间,从而使分类或预测更加容易。与人工规则构造特征方法相比,利用大数据来学习特征更能够刻画数据的丰富内在信息。

朱常宝(1989—),男,硕士生,主要研究方向为图像处理、特征提取、机器学习, E-mail: zhuchangbao1010@163.com;程 勇(1976—),男,讲师,主要研究方向为图像处理、机器学习、模式识别;高 强(1989—),男,硕士生,主要研究方向为图像处理、机器学习。

从本质上来说,深度学习仍然属于一种监督学习模型。通过大量的实例来训练神经元间的权重,可以让整个神经网络按照最大概率来生成训练数据。然而高质量的训练集获取通常代价比较高昂,例如为了对生物肿瘤样本图像进行分类,必须先有一个由专家确认的训练集,这就要求专家为训练集进行辨别和标示,需要耗费大量的人力和财力。而且在标示过程中,由于原始数据存在的错误和不规范,在现实中往往还需要对数据进行预处理。在当今大数据时代,这个预处理工作往往占据了全部项目的多数时间和资源,大大限制了深度学习的实际应用。

我们引入半监督学习的思想来解决上述存在的问题。因为大规模的训练集比较难以获得或不存在,可以对少量的例子进行标示形成一个初始的训练集,然后通过分类或学习过程,充分利用大量的未标示的例子的信息,扩大初始的训练集的规模使之能满足训练深度置信网络的要求。本文中主要研究了3种学习和分类思想来扩展训练集,即Knn、支持向量机(Support Vector Machine, SVM)和pHash算法。我们使用多个受限波尔兹曼机组成的深度置信网络来进行数字的字符特征识别,结果表明半监督深度信念网络比传统多层受限波尔兹曼机有更高的分类准确率。

本文第2节为问题的描述,第3节介绍了半监督深度学习的基本思想和算法,第4节为实验结果的解释和解释,最后总结全文并给出了未来研究的方向。

2 模型定义

本节将会对受限波尔兹曼机^[14]的一些模型变量进行形式化的表示,以使后面的论述更方便、直观。首先定义 $v = (v_1, v_2, \dots, v_i, \dots, v_N) \in (0, 1)^N$ 为模型的输入数据集,其中 v_i 表示第 i 个可见单元的状态,即可视层 v 。针对特定数据集有效标记难获取的情况,通过利用少量数据集来训练深度信念网络模型,通过测试集得到特征提取准确率 $rate_0$ 。然后通过半监督的思想对数据集 v 进行数据集扩充得到 v' ,利用扩充后的数据集进行训练深度信念网络模型,通过测试集得到特征提取准确率 $rate_1$ 。两者的准确率相比较就会得出半监督的思想对于此模型的特征提取有无改善。这里定义隐含层的向量表示为 $h = (h_1, h_2, \dots, h_j, \dots, h_M) \in \{0, 1\}^M$,其中 h_j 表示第 j 个隐单元的状态。

能量函数为:

$$E(v, h | \theta) = \sum_{i=1}^N a_i v_i + \sum_{j=1}^M b_j h_j + \sum_{i=1}^N \sum_{j=1}^M v_i W_{ij} h_j$$

其中, $\theta = \{W_{ij}, a_i, b_j\}$ 是RBM参数,均为实数。其中 W_{ij} 表示可见单元 i 与隐单元 j 之间的连接权重, a_i 表示可见单元 i 的偏置, b_j 表示隐单元 j 的偏置。当参数确定时,得到 (v, h) 联合概率分布:

$$P(v, h | \theta) = \frac{e^{-E(v, h | \theta)}}{Z(\theta)}, Z(\theta) = \sum_{v, h} e^{-E(v, h | \theta)}$$

通过联合概率分布得到边际函数,即似然函数:

$$P(v | \theta) = \frac{1}{Z(\theta)} \sum_h e^{-E(v, h | \theta)}$$

RBM由于层内无连接的特殊结构,具有以下性质:当给

定可见单元的状态时,各隐含单元的激活状态之间是条件独立的。那么,第 j 个隐单元的激活概率为:

$$P(h_j = 1 | v, \theta) = \sigma(b_j + \sum_i v_i W_{ij})$$

其中, $\sigma(x) = \frac{1}{1 + \exp(-x)}$ 为sigmoid激活函数。同理可得到第 i 个可见单元的激活概率为:

$$P(v_i = 1 | h, \theta) = \sigma(a_i + \sum_j W_{ij} h_j)$$

3 基于半监督的深度学习

3.1 半监督

本小节将介绍除了输入输出成对出现的训练样本 $\{(x_i, y_i)\}_{i=1}^n$ 外,通过在学习过程中追加输入训练样本 $\{x_i\}_{i=n+1}^{n+n'}$,来提高学习精度的半监督学习算法。在近年来机器学习算法的具体应用中,上述只作为输入的训练样本的收集变得更加容易了。例如,在网页的分类问题中,为了生成输入输出成对出现的训练样本 $\{(x_i, y_i)\}_{i=1}^n$,需要人们在看到网页 x_i 以后,采用人工输入的方法对其添加标签 y_i 。而对于只有输入的训练样本 $\{x_i\}_{i=n+1}^{n+n'}$ 而言,不要人工介入,互联网本身就可完成数据的收集工作。

对输入输出关系进行监督学习,可以看作是在已知输入 x 时,对输出 y 的条件概率 $p(y|x)$ 进行估计的问题。另一方面,由于只有输入的训练样本中没有输出,只包含输入的概率密度 $p(x)$ 的相关信息。因此,只利用输入的训练样本直接对条件概率 $p(y|x)$ 进行学习,并不会得到很好的效果。半监督学习会首先假定输入概率密度 $p(x)$ 和条件概率密度 $p(y|x)$ 之间具有某种关联,利用对输入概率密度 $p(x)$ 的估计来辅助对条件概率密度 $p(y|x)$ 的估计,进而使得最终的学习精度得以提升。

3.2 算法描述

DBN网络由多个RBM级联而成。上一个RBM的输出就是下一个RBM的输入,这样就组成了一个深度网络。

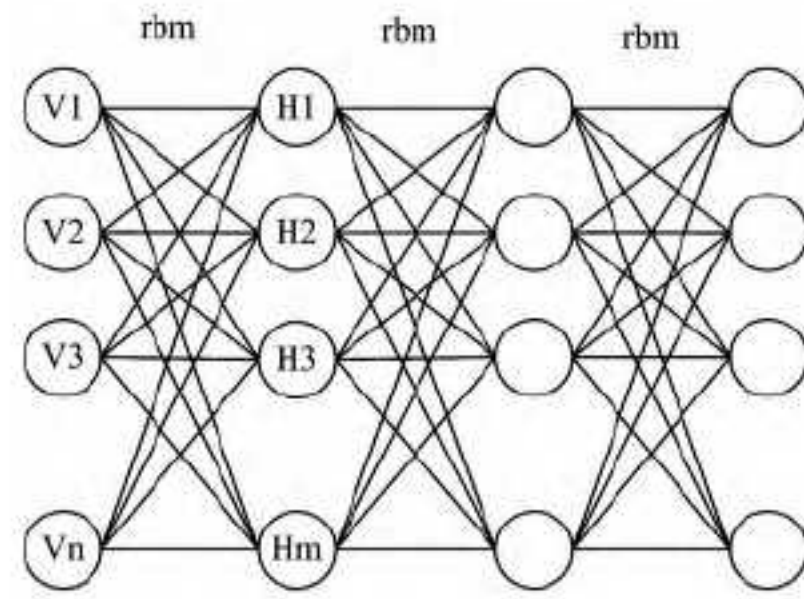


图1 DBN深度网络

图1是一个典型的DBN深度网络模型图,它的标准算法是逐层贪婪算法^[7],而对比散度算法^[8]是训练RBM的标准算法。逐层贪婪算法的具体过程即先利用原始输入来训练网络的第一层,得到其参数 $W^{(1,1)}, b^{(1,1)}, W^{(1,2)}, b^{(1,2)}$;然后网络第一层将原始输入转化成为隐藏单元激活值组成的向量(假设该向量为A),接着把A作为第二层的输入,继续训练得到第二层的参数 $W^{(2,1)}, b^{(2,2)}, W^{(2,1)}, b^{(2,2)}$;最后,对后面的各层采用同样的策略,即将前层的输出作为下一层的输入的方式依次训练。对于上述训练方式,在训练每一层参数时,会固定

其他各层的参数保持不变。所以,如果想得到更好的结果,在上述训练过程完成后,可以通过反向传播算法调整所有层的参数以改善结果,这个过程一般被称为“fine-tuning”^[12]。

已有研究论文中有许多针对该算法的改进,文献^[5]提出了 Hybrid-pretraining 的预训练方法使得 DBN-DNN 的训练初值与全局最佳值相距更近。文献^[6]介绍 Average-SGD 算法,这些方法都可以达到加快网络收敛的目的,从而减少 DBN-DNN 训练中 Fine-tuning 过程的周期数,加快训练速度。文献^[9]提出同等条件下扩大网络的规模可以获得更好的识别性能。文献^[13]提出了采用 SDS 算法来加快训练速度,即在每个 Epoch 中随机地选取部分数据进行 Fine-tuning 训练,并且随着训练的不断深入,选取的数据量呈现逐渐减少的趋势。

本文算法引入了半监督的思想来扩充数据集,虽然数据集的数量增加了,但是会导致其中的一些数据标签特征是错误的,这就会误导 DBN 网络的特征提取的准确性。在使用半监督方法扩充数据集之后,再随机相同比例地划分出训练集和测试集,虽然训练集中存在极少量比例错误的数据,但绝大多数还是有效的,最终提高了 DBN 模型的特征提取的准确率,从而有效地证明了此思想的可行性。

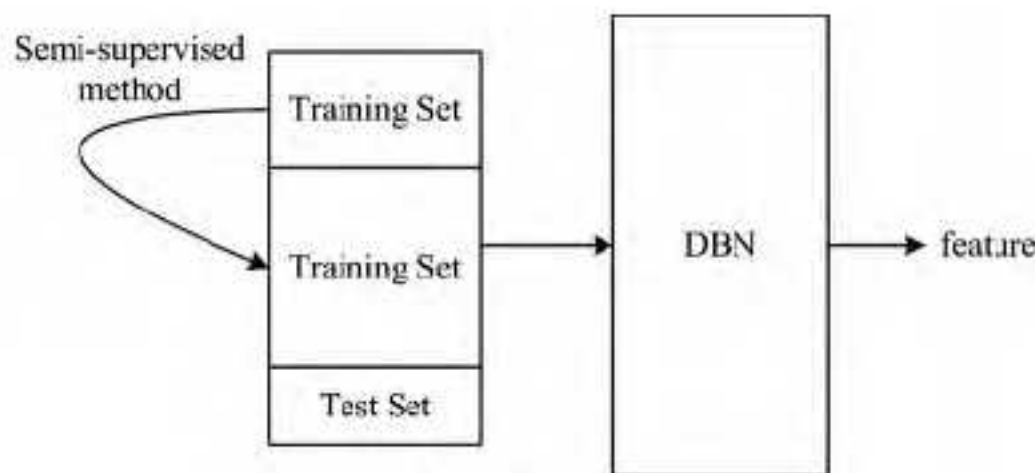


图2 半监督方法

本文方法的流程如图2所示,通过半监督的方法来扩充原有的数据,进而学习训练 DBN 网络模型。算法代码如下:

半监督学习,运用少量数据集 X_0 扩充数据集,标记数据标签组成新的数据集 X_1
 输入:扩充后的数据集 X_1 动量 m , 权衰减 l_0 , 学习率 ϵ
 输出: 链接矩阵 W 、可见层偏置向量 a 、隐含层偏置向量 b
 训练:
 初始化: 令各个 RBM 可见层单元的初始状态 $v_1 = x_1$; W 为随机的标准分布, a, b 为 0
 For $n=1, 2, \dots, N$ (有 N 个 RBM)
 For $t=1, 2, \dots, T$ (迭代 T 次)
 For $j=1, 2, \dots, m$ (隐单元)
 计算 $P(h_{1j}=1|v_1)$, 即 $P(h_{1j}=1|v_1) = \sigma(b_j + \sum_i v_{1i} W_{ij})$
 EndFor
 For $i=1, 2, \dots, n$ (可见单元)
 计算 $P(v_{2i}=1|h_1)$, 即 $P(v_{2i}=1|h_1) = \sigma(a_i + \sum_j W_{ij} h_{1j})$
 EndFor
 For $j=1, 2, \dots, m$ (隐单元)
 计算 $P(h_{2j}=1|v_2)$, 即 $P(h_{2j}=1|v_2) = \sigma(b_j + \sum_i v_{2i} W_{ij})$
 EndFor
 迭代参数
 $-W \leftarrow mW + \epsilon((P(h_1=1|v_1)v_1^T - P(h_2=1|v_2)v_2^T) - l_0 W)$
 $-a \leftarrow ma + \epsilon(v_1 - v_2)$

$$-b \leftarrow mb + \epsilon(P(h_1=1|v_1) - P(h_2=1|v_2))$$

EndFor

EndFor

3.3 算法分析

对于上面提到的采用半监督方法来扩充数据集,本文用到了 Knn、SVM、pHash 算法。下面就逐一介绍这 3 种算法。

Knn 算法是一种理论上比较成熟的方法,也是最简单的机器学习算法之一。该方法的思路是:如果一个样本在特征空间中的 k 个最相似(即特征空间中最邻近)的样本中的大多数属于某一个类别,则该样本也属于这个类别。Knn 算法中,所选择的邻居都是已经正确分类的对象。该方法在定类决策上只依据最邻近的一个或者几个样本的类别来决定待分样本所属的类别。Knn 方法虽然在原理上也依赖于极限定理,但在类别决策时,只与极少量的相邻样本有关。由于 Knn 方法主要靠周围有限的邻近的样本,而不是靠判别类域的方法来确定所属类别的,因此对于类域的交叉或重叠较多的待分样本集来说, Knn 方法较其他方法更为适合。

SVM 算法^[10]通过一个非线性映射 p , 把样本空间映射到一个高维乃至无穷维的特征空间中(Hilbert 空间),使得在原来的样本空间中非线性可分的问题转化为在特征空间中的线性可分的问题。简单地说,就是升维和线性化。升维,就是把样本向高维空间做映射,一般情况下这会增加计算的复杂性,甚至会引起“维数灾难”,因而人们很少问津。但是作为分类、回归等问题来说,很可能在低维样本空间无法线性处理的样本集,在高维特征空间中却可以通过一个线性超平面实现线性划分(或回归)。一般的升维都会带来计算的复杂化, SVM 方法巧妙地解决了这个难题,应用核函数的展开定理,就不需要知道非线性映射的显式表达式。由于是在高维特征空间中建立线性学习机,因此与线性模型相比,不但几乎不增加计算的复杂性,而且在某种程度上避免了“维数灾难”,这一切要归功于核函数的展开和计算理论。

pHash 算法^[11]是均值哈希的增强版,均值哈希虽然简单,但是受均值的影响非常大。例如对图像进行伽马校正或者直方图均衡就会影响均值,从而影响最终的 hash 值。pHash 将均值的方法发挥到了极致。使用离散余弦变换来获取图片的低频成分。

离散余弦变换是一种图像压缩方法,即将图像从像素域变换到频率域。一般图像都存在很多冗余和相关性,所以变换到频率域之后,只有很少的一部分频率分量的系数不为 0,大部分的系数都为 0。

pHash 算法如下:

- (1) 缩小尺寸: pHash 从缩小图片开始,但图片大于 8×8 , 32×32 是最好的。其目的是简化离散余弦变换的计算,而不是减小频率。
- (2) 简化色彩: 将图片转化为灰度图像,进一步简化计算量。
- (3) 计算离散余弦变换: 计算图片的离散余弦变换,得到 32×32 的离散余弦变换系数矩阵。
- (4) 缩小离散余弦变换: 虽然离散余弦变换的结果是

32 * 32 大小的矩阵,但只要保留左上角的 8 * 8 的矩阵,这部分呈现了图片中的最低频率。

(5) 计算平均值,计算所有 64 个像素的离散余弦变换之后矩阵的平均值。

(6) 计算 hash 值,这是最主要的一步,根据 8 * 8 的离散余弦变换矩阵,设置 0 或 1 的 64 位的 hash 值,大于或等于离散余弦变换均值的设为“1”,小于均值的设为“0”。组合在一起,就构成了一个 64 位的整数,这就是一张图片的指纹。组合次序不重要,只要保证所有图片都采用相同的次序即可。

如果图片放大或缩小,或改变纵横比,结果值也不会改变。增加或减少亮度、对比度,或改变颜色,对 hash 值都不会有太大的影响。PHash 算法的最大优点是计算效率较高,运算时间较短,并且能够达到较好的最终结果。

比较图片相似性的方法为先计算出两张图片的 hash 指纹,即 64 位 0 或 1 值,然后计算不同位的个数(汉明距离)。如果汉明距离为 0,表示两张图片非常相似;如果汉明距离小于 5,则表示有些不同,但比较相近;如果汉明距离大于 10,则表明两个图片是完全不同的。

4 实验

本实验使用的是 MNIST 字符集,这个数据集共包含 70000 幅图像(60000 幅训练图像和 10000 幅测试图像)。每幅 MNIST 图像是单一的手写的数字字符的数字化图片,每幅图像的大小为 28 * 28 像素。像素值为 0,表示白色;为 255,表示黑色;中间像素值表示灰度级。

本文所选用的实验方案有 3 个,第一个是监督学习,选取 120 幅训练图像,训练 DBN 网络,再用 30 幅测试图像来测试 DBN 网络。第二个加入半监督的思想扩充训练集,选用 Knn 和 SVM 算法基于 120 幅训练图像来扩充训练集到 300 幅图像,接着选取 100 幅测试图像来测试。第三个也是选用 pHash 算法来进行半监督的扩充训练集,接着也是进行测试。通过对比监督学习和半监督学习的测试准确率来说明半监督思想的有效性。

选用的 DBN 网络共由 3 个 RBM 连接组成,第一个 RBM 为 784 * 500,第二个 RBM 为 500 * 500,第三个 RBM 为 510 * 2000,其中第三个 RBM 在第二个 RBM 输出的 500 个隐含单元值的基础上加上 10 个标签标记数组。

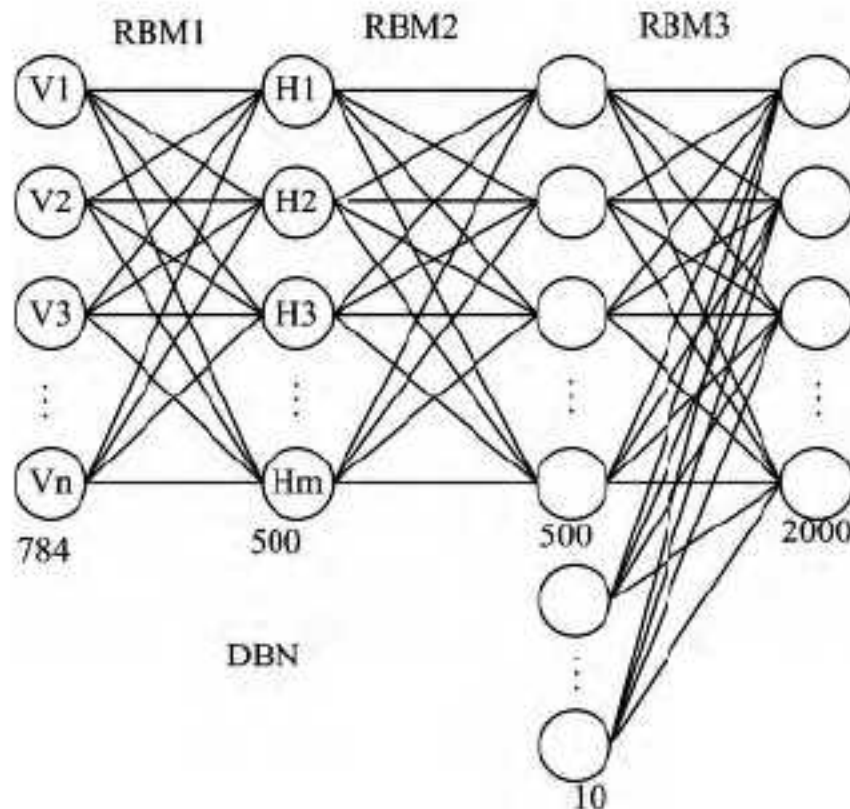


图 3 DBN 模型图

DBN 模型图如图 3 所示,它在最后一个 RBM 加入了特征提取的神经元,当输入的 MNIST 数据的标签为几时,那么对应的 501—510 单元就会有对应的单元为 1,其他的为 0。

实验结果:

首先选取了 150 个数据集,其中 120 个为训练集,30 个为测试集。因为传统的深度学习是监督的,基于其进行了实验。其中所选的参数如下:学习率为 0.1,权衰减为 0.0002,学习次数小于 10 时动量为 0.5,大于 10 时其为 0.9。

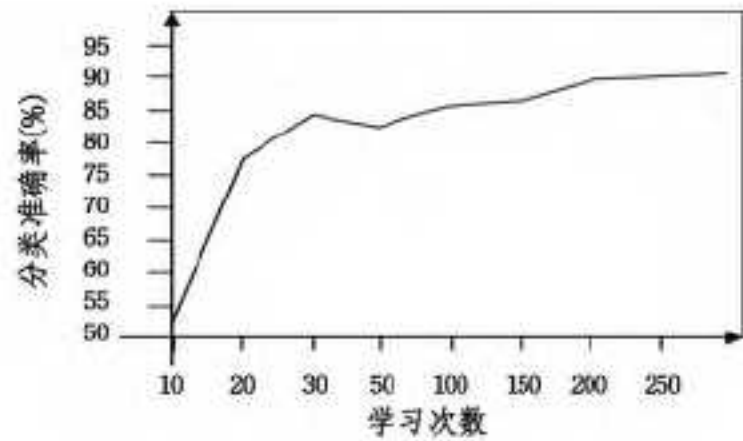


图 4 监督学习数据

图 4 为随着迭代次数的增加监督学习中得到的特征提取的分类准确率基本走势图。具体的数据如表 1 所列。

表 1 监督学习模型数据

迭代次数	10	20	30	50	100	150	200
测试准确数	15.798	23.199	25.398	24.798	25.599	25.794	26.598
准确率(%)	52.66	77.33	84.66	82.66	85.33	85.98	88.66

先通过半监督的算法进行数据集的扩充,从 150 个数据扩充到 400,其中学习了 250 个数据。

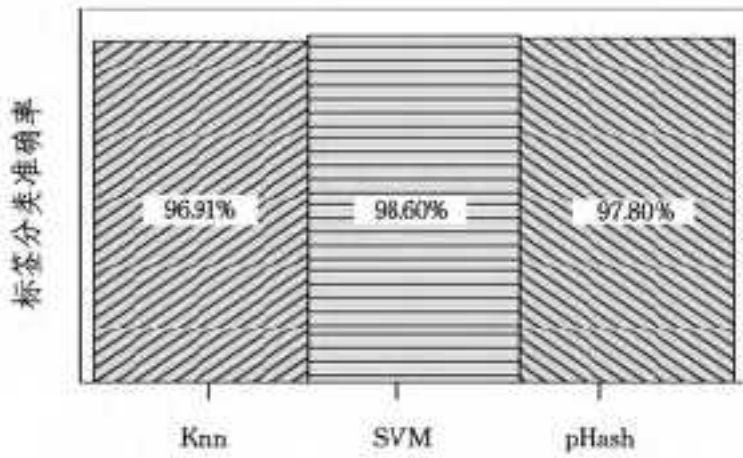


图 5 半监督扩充数据集

图 5 展示了 3 种扩充数据集的方法得到标签的准确率,由三者的表现结果来看,得到的数据都有一定的错误,这些数据也会一同训练 DBN 模型,具体的数据如表 2 所列。

表 2 扩充数据集数据

方法	标签分类准确率 (%)	正确标签分类个数
Knn	96.91	242.257
SVM	98.60	246.5
pHash	97.80	244.5

通过扩充后总数为 400 的数据集来学习 DBN 网络,其中训练集为 300,测试集为 100。虽然测试集中有错误的标签,但是本实验通过原始的正确标签来比对最终的结果,如图 6 所示。

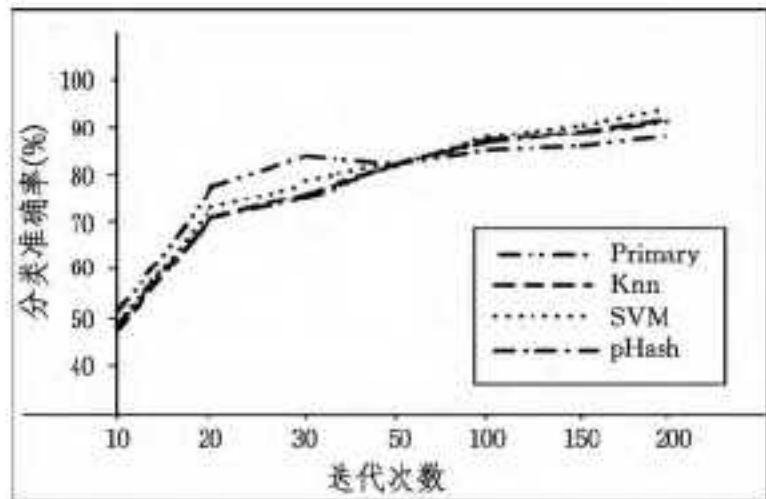


图 6 各方法对比

图6中通过迭代次数的增加对比了监督和半监督,共有4条线。基本上最终结果都呈现出良好的走势,但是在扩充之后可以稍微地改善最终的特征提取。采用各种方法得到的具体数据如表3所列。

表3 3种方法迭代数据(%)

迭代次数	10	20	30	50	100	150	200
Knn	47.66	71.35	76.42	82.42	87.41	89.14	91.76
SVM	49.52	73.16	78.45	83.22	87.82	90.46	93.45
pHash	47.25	71.56	75.85	82.10	86.45	88.45	91.42

通过本实验得出基于半监督的思想提高了特征提取的准确率,在多次试验之后得到的数据不是整数,而是平均之后的数据。经过实验数据的比对得出,随着迭代次数的不断提高,标签分类的准确率是逐步提高的,但并不是线性的走势。其中相比于少量的数据集训练DBN模型的测试准确率,基于半监督的思想扩充的数据集训练模型得到的测试准确率提高了2%~5%。如果想得到有效的提高,必须使在初始数据扩充数据集上的准确率保持在95%以上。此实验模拟了现实经常存在的现象,有效数据集中的个数比较少,但是通过人工进行前期的特征提取既不便又耗费人力物力,但是可以通过半监督的思想来扩充数据集,进而提高DBN模型的特征提取准确率。在本文中只是有效地证明了此思想的有效性,但是针对不同的问题,仍然需要各自不同的方法来进行不同数据集的模型训练。

结束语 目前,深度学习在语音识别、图像分类、中文文本分类、机器学习等领域中成为“标准”技术,但其效果和性能严重依赖训练集和训练过程。本文主要研究半监督学习思想来扩充数据集以提高深度学习模型的性能。我们选取Knn, SVM和pHash算法来扩充训练集,并对算法的学习结果进行了分析和比较。通过对比监督学习算法同样训练DBN网络,可以总结出半监督的学习类别标签的准确率达到95%以上后即可提高DBN的测试准确率。相对于监督的DBN网络训练,半监督的算法平均提高了3%的准确率。在3种标示算法中,又以SVM更优。但是训练DBN特征的算法中,由于选取的参数不同,可能得出不同的实验结果。在训练时间上还是比较快的,毕竟选取的训练图像和测试图像比较少。以后的工作重点为研究如何通过更优的半监督算法来提高DBN的识别准确率。

参考文献

[1] Hinton G, Deng L, Yu D, et al. Deep neural networks for acoustic modeling in speech recognition; The shared views of four re-

search groups[J]. Signal Processing Magazine, 2012, 29(6): 82-97

- [2] Mohamed A, Dahl G E, Hinton G. Acoustic modeling using deep belief networks[J]. IEEE Trans on Audio, Speech, and Language Processing, 2012, 20(1): 14-22
- [3] Roux N L, Bengio Y. Representational power of restricted Boltzmann machines and deep belief networks[J]. Neural Computation, 2006, 20(6): 1631-1649
- [4] Hinton G E, Osindero S, The Y. A fast learning algorithm for deep belief nets[J]. Neural Computation, 2006, 18(7): 1527-1554
- [5] Sainath T N, Kingsbury B, Soltan H, et al. Optimization technique to improve training speed of deep belief networks for large speech tasks[J]. IEEE Trans on Audio, Speech, and Language Processing, 2013, 21(11): 2267-2276
- [6] You Z, Wang X, Xu B. Exploring one pass learning for deep neural network training with averaged stochastic gradient descent [C]// IEEE International Conference on Acoustics, Speech and Signal Processing. 2014: 6854-6858
- [7] Erhan D, Courville A, Bengio Y, et al. Visualizing Higher Layer Features of a Deep Network [C]// Spotlight presentation and poster at the ICML 2009 Workshop on Learning Feature Hierarchies. Montreal, Canada, 2009
- [8] Hinton G E. Training products of experts by minimizing contrastive divergence[J]. Neural Computation, 2002, 14(8): 1771-1800
- [9] Mass A L, Hannum A Y, Lengerich C T, et al. Increasing deep neural network acoustic model size for large vocabulary continuous speech recognition[J]. arXiv preprint arXiv: 1406. 7806, 2014
- [10] Hsu C, Lin C. A comparison on methods for multi-class support vector machines[J]. IEEE Transactions on Neural Networks, 2002, 12: 415-425
- [11] Zauner. Christoph: Implementation and Benchmarking of Perceptual Image Hash Functions [C]// Master's thesis, Upper Austria University of Applied Sciences. Hagenberg Campus, 2010
- [12] Hinton G E. A practical guide to training restricted Boltzmann machines[J]. Momentum, 2010, 9(1): 926
- [13] Deng W, Qian Y, Fan Y, et al. Stochastic data sweeping for fast DNN training [C]// Proc of IEEE International Conference on Acoustics, Speech and Signal Processing. 2014: 240-244
- [14] 罗恒. 基于协同过滤视角的受限波尔兹曼机研究 [D]. 上海, 上海交通大学, 2011