

结合空间—光谱信息的快速自训练高光谱遥感影像分类

金垚^{1,2,3}, 董燕妮^{1,2,4}, 杜博⁵

1. 中国地质大学(武汉) 地球物理与空间信息学院, 武汉 430074;

2. 湖北珞珈实验室, 武汉 430079;

3. 武汉大学 测绘遥感信息工程国家重点实验室, 武汉 430079;

4. 武汉大学 资源与环境科学学院, 武汉 430079;

5. 武汉大学 计算机学院, 武汉 430072

摘要: 自训练方法被广泛应用于高光谱影像分类任务中以解决标记样本获取困难的问题。传统的自训练方法不仅忽略了高光谱影像所能提供的空间信息, 导致最终分类精度受到影响; 同时在每次迭代过程中都需要完成一次对未标记数据的分类任务, 导致需要大量的时间成本。因此, 针对上述问题, 本文提出了一种基于空间—光谱信息的快速自训练方法用于高光谱影像分类。与传统的自训练方法不同, 该方法在迭代过程中使用空间—光谱信息对未标记数据进行筛选完成标记样本的扩充, 而不是使用分类器对未标记样本进行分类。首先针对初始标记样本使用空间邻域块选择空间近邻点, 然后使用自适应阈值对空间近邻点进行二次筛选得到空谱近邻点赋予标记, 最后根据扩充后的标记样本对分类器进行训练完成分类任务。结果表明, 在 Washington DC Mall Subimage 高光谱数据集中每类分别选择 2 个和 10 个训练样本时, 整体分类精度分别达到了 93.17% 和 95.43%; 而在 Indian Pines 数据集中整体分类精度分别达到了 59.75% 和 86.13%。我们提出的结合空间—光谱信息的快速自训练方法和对比方法相比, 我们的方法有明显的提升。

关键词: 高光谱遥感, 半监督分类, 小样本问题, 空间—光谱信息, 自训练方法

中图分类号: P237/P2

引用格式: 金垚, 董燕妮, 杜博. 2024. 结合空间—光谱信息的快速自训练高光谱遥感影像分类. 遥感学报, 28(1): 219–230

Jin Y, Dong Y N and Du B. 2024. Fast self-training based on spatial-spectral information for hyperspectral image classification. National Remote Sensing Bulletin, 28(1): 219–230 [DOI: 10.11834/jrs.20232286]

1 引言

高光谱遥感光谱分辨率高, 波段数目众多, 具有图谱合一的特性。这些特点使得高光谱遥感影像拥有更为丰富的光谱波段信息, 能够实现地物精准分类, 被广泛应用于精准农业、生态监测、土地利用分类等领域 (Dong 等, 2022; Plaza 等, 2009; Jin 等, 2022)。

传统的分类算法大多数为监督分类算法, 即通过标记样本提供的先验信息完成对分类器的训练, 然后使用经过训练的分类器对未标记样本进行分类 (Blaschke, 2010)。在监督分类算法中,

根据有无空间信息的参与, 可以进一步分为基于光谱信息的分类算法和基于空—谱信息的分类算法。基于光谱信息的分类算法仅仅依靠光谱信息对高光谱图像进行分类, 其中最为经典的算法就是支持向量机 SVM (Support Vector Machines) (Zhang 等, 2004) 和多层感知机 MLP (Multilayer Perceptron) (Zerguine 等, 2001)。SVM 以结构风险最小化为准则, 寻找最佳的分类面实现分类; MLP 则以神经网络的结构来对输入数据进行高效分类。而基于空—谱信息的分类算法在利用光谱信息的同时, 还利用了高光谱图像中蕴含的空间信息。其中基于组合核 CK (Composite Kernel)

收稿日期: 2022-06-06; 预印本: 2022-12-24

基金项目: 国家自然科学基金(编号: 62222116, 62171417, 41871243, 62141112); 湖北珞珈实验室开放基金(编号: 220100058)

第一作者简介: 金垚, 研究方向为高光谱遥感影像降维与分类。E-mail: yao.jin@whu.edu.cn

通信作者简介: 董燕妮, 研究方向为遥感图像智能解译。E-mail: dongyanni@cug.edu.cn

(Camps-Valls 等, 2006) 的方法是最基础的空—谱分类方法, CK 采用空间窗口提取窗口中心样本点的邻域信息, 得到空间信息核, 然后与光谱信息核相结合, 得到 CK 替换 SVM 中的核函数, 有效提高了高光谱图像的分类精度。

然而, 在监督分类算法中训练 1 个较优的分类器需要大量标记数据的支撑, 尤其是高光谱遥感数据, 大量的光谱波段意味着需要更多的先验信息 (Li 等, 2018; Ghamisi 等, 2017)。但是由于各种限制因素, 收集标记样本的成本是十分昂贵的, 在实际应用中往往无法获得足够的标记样本来支撑分类器的训练 (Bioucas-Dias 等, 2013)。

近些年, 已有研究通过半监督分类算法来解决标记样本不足的问题 (van Engelen 和 Hoos, 2020)。半监督分类算法是在只有少量标记数据的条件下, 结合未标记数据中蕴含的信息辅助分类器的训练, 能够有效增加分类器的泛化性和准确率 (Camps-Valls 等, 2007)。一般而言, 半监督分类算法可以分为: 协同训练方法 (Li 等, 2020)、基于图的方法 (Chen 等, 2017) 以及自训练方法 (Ge 等, 2021)。

协同训练方法在不同视图上训练得到多个分类器, 通过多个分类器之间的协同作用实现对无标记样本的分类。Tri-Training (Zhou 和 Li, 2005) 将标记数据分为 3 个视图, 根据标记数据训练出对应的 3 个分类器, 然后采用集成学习对未标记样本进行标注; TT-AL-MSH (Tri-Training with Active Learning and Multi-Scale Homogeneity) (Tan 等, 2016) 基于多样性的考量, 通过在 4 个分类器中选择 3 个具有相当大差异的分类器, 改进了传统的 Tri-Training, 有效的提高了分类器的分类能力; Multi-Train (Gu 和 Jin, 2017) 将 Tri-Training 进一步扩展到更多的视图以提高性能。但是协同训练要求多视图之间存在独立性, 在实际应用时往往难以满足 (Wang 等, 2014)。

基于视图的方法认为在特征域中属于同一邻域的样本有更大的概率属于同一类别。LP (Label Propagation) (Wang 和 Zhang, 2008) 根据所有样本之间的相似性构建 1 个权重图, 然后各个样本在其相邻的样本间进行标签传播; NLP (Negative Label Propagation) (Zoidi 等, 2018) 通过负标签信息有效的指导半监督学习的过程; SaGSSL (Safety-Aware Graph-Based Semi-Supervised Learning) (Gan

等, 2018) 通过样本余量来判断所构建的图是否合适, 能够自适应地选择合适的图, 并学习 1 个安全的半监督分类器; 但是基于图的方法对图构建有很高的计算要求, 需要较高的时间成本 (Ge 等, 2021)。

自训练方法作为半监督学习中的经典方法, 具有简单以及不需要先验假设的优点 (van Engelen 和 Hoos, 2020)。自训练方法首先利用标记数据训练 1 个初始的分类器对未标记数据进行分类, 然后选择置信度比较高的样本进行标记, 最后使用扩充后的标记数据集完成对分类器的训练。在整个自训练过程中, 最为关键的部分就是如何获得置信度高的分类结果, 因为初始的分类器是在少量标记数据上进行的训练, 训练样本的数量不足, 会导致分类器训练不充分, 容易对未标记数据产生误判。而误标记样本产生的错误在迭代过程不断累加, 进而导致最终的分类器分类效果较差 (Gu, 2020)。

因此为了提高自训练方法的性能, 已有研究提出不同自训练方法的改进版本。TSVM (Transductive Support Vector Machines) (Bruzzone 等, 2006) 采用动态阈值为基于自训练方法的 SVM 分类器选择有信息的样本; SBSL (Segmentation-Based Self-Learning) (Lu 等, 2016) 使用高分辨影像的分割技术来选择信息量大且多样化的样本, 以便在自训练中有效地训练监督分类器; SP-SVM (Self-Training Approach with Pseudo) (Ayday 和 Minz, 2018) 使用伪标签验证集提升模型分类精度, 减少误差; ST-HP (Self-Training Hierarchical Prototype-Based) (Gu, 2020) 利用数据集中存在的层次结构, 以自训练的方式训练了 1 个基于层次结构的分类器, 以改善遥感图像的半监督分类。但是以上的自训练方法通常存在着 2 个问题: (1) 忽略了高光谱影像提供的空间信息, 导致最终分类精度受到影响; (2) 每轮迭代都需要使用分类器对未标记样本进行分类, 需要大量的时间成本。

针对上述问题, 本文提出 1 种基于空间—光谱信息的快速自训练方法 FST-SS (Fast Self-Training with Spatial-Spectral Information)。FST-SS 在每次迭代中首先利用空间信息为标记样本提取空间近邻点, 然后根据标记样本的光谱邻域信息计算自适应阈值对空间近邻点进行筛选得到标记样本的空谱近邻点, 最后将空谱近邻点视为标记样本同

类点赋予同样的标签并加入到标记样本集中进行下1次迭代, 在训练过程中避免了每次迭代所需的分类过程, 大大降低了时间成本。本文采用 Washington DC Mall Subimage 和 Indian Pines 2 个常用的高光谱数据集, 对比 1 种常用的监督分类算法以及 4 种半监督分类算法来验证方法的有效性。

2 光谱信息快速自训练方法

对于高光谱影像, 本文使用 $X = [x_1, x_2, \dots, x_n] \in R^{d \times n}$ 代表标记样本集, 其中 d 表示影像波段数, n 表示标记样本的个数。 $Y = [y_1, y_2, \dots, y_n] \in R^{n \times 1}$ 表示对应的标签信息。根据已知的标

签信息, 可以将标记样本集分为两部分。如果 $y_i = y_j$, 标记样本 x_i 和 x_j 属于相似样本对集合 S , 反之则属于不相似样本对 D 。

$$S: \forall (x_i, x_j) |_{y_i = y_j} \in S, \quad (1)$$

$$D: \forall (x_i, x_j) |_{y_i \neq y_j} \in D. \quad (2)$$

针对传统自训练方法中缺乏对空间信息的考量以及计算效率较低的问题, FST-SS 引入空间信息, 通过使用空间—光谱信息直接对未标记数据进行筛选, 避免每轮迭代中的分类过程, 能够在保证最终精度的同时, 极大的提高了计算效率, FST-SS 具体流程如图 1 所示。

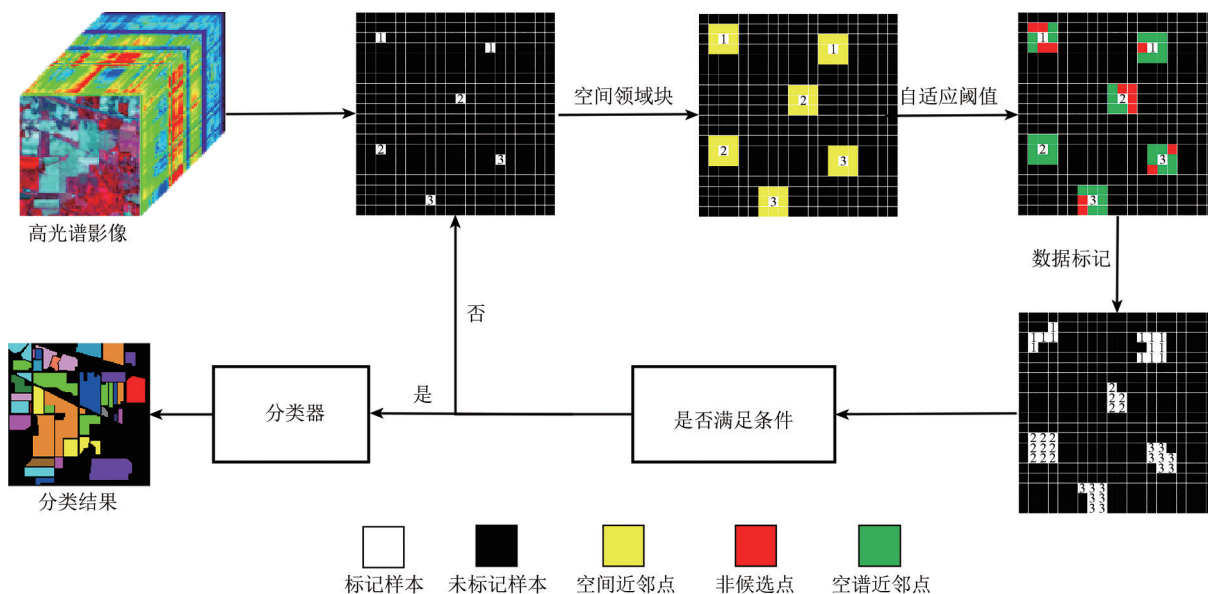


图1 FST-SS 流程图

Fig.1 FST-SS flow chart

2.1 空间近邻点选择

根据地学第一定律可知, 地物之间存在着空间相关性, 所以高光谱影像中像素和地物的空间分布在一定范围内是一致的。因此空间上相邻的像素有极大概率是同 1 种地物。本文遵循这种假设, 使用 3×3 的空间邻域块为标记样本选择空间近邻点。标记样本处于块的中心, 其余的点则为标记样本点的空间近邻点。通过这样一种方式每个标记样本可以最多选择 8 个空间近邻点, 而当标记样本位于影像边缘时, 空间近邻点则少于 8 个。

2.2 利用自适应阈值筛选空谱近邻点

尽管标记样本与周围的空间近邻点有着极高的相似度, 但是无法保证每 1 个空间邻近点都与中

心的标记样本点类别相同。因此本文引入阈值 b 对空间近邻点进一步筛选, 筛选出空谱近邻点。如果 $d(x_i, x_j) > b$, 则认为 x_i, x_j 属于不同类点, 如果 $d(x_i, x_j) < b$, 则认为 x_i, x_j 属于同一类点。 $d(x_i, x_j)$ 为 x_i, x_j 之间的光谱距离, 计算公式如下:

$$d(x_i, x_j) = \|x_i - x_j\|_2 \quad (3)$$

考虑到每 1 个标记样本的光谱邻域信息并不相同, 当面对复杂的数据时, 采用固定的阈值缺乏对邻域信息的考虑, 往往难以取得较好的结果。所以为了充分考虑到每 1 个标记样本的光谱邻域信息, 做到具体问题具体分析, 本文采用 1 种自适应阈值, 根据每 1 个标记样本的光谱邻域信息计算相对应的阈值, 用于筛选空谱近邻点。具体设置如

下：对于每一个标记样本 $\mathbf{x}_i$ ，首先根据式 (3) 计算 $\mathbf{x}_i$ 与其同类标记样本之间的光谱欧式距离，得到同类标记样本之间的距离矩阵，将距离 $\mathbf{x}_i$ 最近点的距离作为阈值，即

$$b_i = \min d(\mathbf{x}_i, \mathbf{x}_{ik}), \mathbf{x}_i, \mathbf{x}_{ik} \in S \tag{4}$$

随着迭代次数的不断增加，每次阈值 b_i 的值也会快速减少，被判定为空谱近邻点的要求也会严苛，这样保证了新增标记样本的准确度。但是随着阈值 b_i 不断下降，也有一部分潜在的同类样本没有被纳入到新增标记数据集中去。因此本文还在 b_i 的基础上，考虑到每次迭代 b_i 的缩减，设计了另外 1 种阈值 $b_{it} = \min_t d(\mathbf{x}_i, \mathbf{x}_{ik}), \mathbf{x}_i, \mathbf{x}_{ik} \in S$ ，代表 $\mathbf{x}_i$ 与其第 t 近的同类标记样本之间距离， t 的值为迭代的次数。这样可以减缓 b_{it} 的衰减速度，可以把更多的未标记样本纳入标注数据集中。

2.3 迭代训练

通过上述两步可以完成 1 次迭代的训练过程，将训练得到的新增标记样本集与初始标记样本集合并后，一起投入下 1 次的迭代训练中，直至标记数据集数目不在增加或者达到设定的最大迭代次数。由于本文中产生标记数据的方法不是采用传统的分类方法，并不是所有的数据都会被划分进标记数据集，所以当最后的迭代完成后，仍需要使用分类器对结果进行分类。本文采用 1NN (1-Nearest Neighbor) 分类器，当迭代训练完成后，使用扩充后的标记数据集对分类器进行训练，然后对测试数据完成分类。

3 数据结果处理与分析

3.1 实验数据

3.1.1 Washington DC Mall Subimage 数据集：该数据集由为 HYDICE 传感器于华盛顿广场获得的子影像区域。影像包括 210 个光谱带。波长范围从 0.4—2.4 μm ，影像大小为 280×307 像素。在删除了低信噪比和水吸收带后，总共剩下 191 个波段用于本次实验。图 2 (a) 和 (b) 分别显示了对应的假彩色图像和真实地物图。Washington DC Mall Subimage 数据集在 8 类地物中共有 14266 个样本点，表 1 中列出了每一类的样本数量。

3.1.2 Indian Pines 数据集：由 NASA 的机载可

见光/红外成像光谱仪 (AVIRIS) 于 1996 年在印第安纳州西北部农业区获得。这个场景包括 220 个光谱带，波长范围从 0.4—2.5 μm ，影像大小为 145×145 像素。在删除了低信噪比和水吸收带后，总共剩下 200 个通道用于本次实验。图 3 (a) 和 (b) 分别显示了对应的假彩色图像和真实地面图。该数据集在 16 类特征中共有 10249 个样本点，表 2 中列出了每 1 类的样本数量。

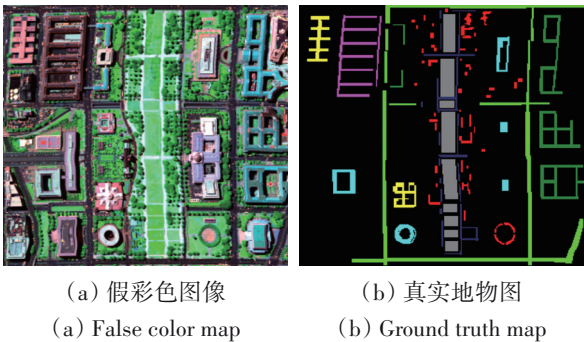


图 2 Washington DC Mall Subimage 数据集
Fig.2 Washington DC Mall Subimage dataset

表 1 Washington DC Mall Subimage 数据真实地物样本信息
Table 1 Ground truth of main classes on Washington DC Mall Subimage

标签	类别	样本数量
1	草地	3133
2	道路	4146
3	房屋 1	1067
4	房屋 2	1928
5	阴影	1013
6	房屋 3	818
7	小路	1096
8	树木	1065
总计		14266

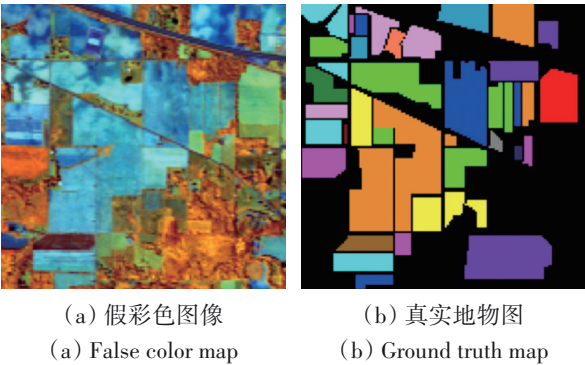


图 3 Indian Pines 数据集
Fig.3 Indian Pines dataset

表2 Indian Pines数据真实地物样本信息
Table 2 Ground truth of main classes on Indian Pines

标签	类别	样本数量
1	苜蓿	46
2	免耕玉米地	1428
3	少耕玉米地	830
4	玉米	237
5	草地/牧场	483
6	草地/树木	730
7	修建的草地	28
8	甘草—堆料	478
9	燕麦	20
10	免耕大豆地	972
11	少耕大豆地	2455
12	大豆	593
13	小麦	205
14	树木	1265
15	建筑	386
16	石钢塔	93
总计		10249

3.2 实验设置

为了证明 FST-SS 的有效性，本文将 FST-SS 与监督分类算法 1NN、半监督分类算法 Tri-Training、Star-SVM (Self-Training with Adaptive Regularization Support Vector Machine) (Cheung 和 Li, 2017)、ST-DP (Self-training based on density peaks of data) (Wu 等, 2018) 以及 LeMA (Learnable Manifold Alignment) (Hong 等, 2019) 进行对比实验，其中 FST-SS-fixed 表示使用阈值 b_i 、FST-SS-unfixed 表示使用阈值 b_u ，所有算法均使用相同的训练样本进行实验，以保证对比的公平性。对于每 1 种数据集，本文在每类地物中随机选择 2、10 样本作为训练样本，余下数据作为测试样本。

图 4 是对最大迭代次数的参数分析结果图。可见：FST-SS 在迭代次数为 50 时结果已经趋于稳定；迭代次数未达到 50 次时就停止了训练。因此本文设置最大迭代次数为 50。此外为了减少随机取样带来的误差，每组实验独立进行 5 次，最终结果为 5 次实验的平均值。实验中采用整体分类精度 OA 以及 kappa 作为定量评价指标，结果中加粗结果为最好结果。

3.3 实验结果

3.3.1 Washington DC Mall Subimage 数据集

表 3 和表 4 分别表示是每类随机选择 2 和 10 个

训练样本的定量评价结果，从结果中可以发现，本文提出的 FST-SS-unfixed、FST-SS-fixed 拥有最优和次优的精度和 kappa 值，而且在大多数地物分类中都能取得优异的分类精度。在每类随机选择 2 个训练样本的实验组中，4 种对比的半监督分类算法的精度相较于 1NN 分类器都有不错的提升。所提出的算法分别取得最优和次优的精度，其中 FST-SS-fixed 精度的提升较小，而 FST-SS-unfixed 提升最多可达 16%，这也证明了 FST-SS 算法在小样本条件下的适用性。而随着训练样本的增多，在每类选择 10 个训练样本的实验组中，提出的算法同样保持着最优和次优的精度，在训练样本增多的情况下 FST-SS-fixed 的精度大幅度增加，远远超出其他的对比算法，而 FST-SS-unfixed 则依旧保持最高的精度。图 5 和图 6 分别为 Washington DC Mall Subimage 数据集上不同初始训练样本的分类结果。可见：在每类选择 2 个训练样本时，本文提出的算法在第 6 类房屋 3（黄色地块）中有着最优的分类结果，其分类准确度远超其他对比算法，而 FST-SS-unfixed 相较于 FST-SS-fixed 而言，在第 5 类阴影（紫色地块）中拥有更平滑的分类结果；在每类选择 10 个训练样本时，本文提出的算法在第 3 类房屋 1（浅蓝色地块）拥有更好的分类结果。图 7 为不同标记样本下的整体分类精度 OA 值。可见分类精度会随着初始标记样本数的增加而增加，当每类标记样本数到达 20 时 OA 值趋于稳定。

3.3.2 Indian Pines 数据集

表 5 和表 6 分别为每类随机选择 2、10 个训练样本在 Indian Pines 数据集上的定量评价结果。可见：在每类随机选择 2 个训练样本的实验组中，提出的算法分别取得最优和次优的精度，其中 FST-SS-fixed 精度的提升达到 18%，而 FST-SS-unfixed 提升最多可达 27%；在每类选择 10 个训练样本的实验组中，提出的算法同样保持着最优和次优的精度。图 8 和图 9 分别为 Indian Pines 数据集上不同训练样本的分类结果图。可见由于始训练样本数目少，所以整体精度较低，分类结果图中噪声较多；本文提出的 FST-SS 拥有更为平滑的分类结果，例如在第 8 类甘草—堆料（红色地块）和第 15 类建筑（浅蓝色地块）中可以明显的看出不同地块的边界，整体的噪声比较少。图 10 展示了不同标记样本下的整体分类精度 OA 图，由于 Indian

Pines 数据集样本分布不均匀，部分类别总体样本较少，所以第 1、7、9 类选择初始标记样本时，选择 50% 作为初始标记样本。得到的结果与

Washington DC Mall Subimage 数据集一致，分类精度会随着初始标记样本数的增加而增加，当每类标记样本数到达 40 时趋于稳定。

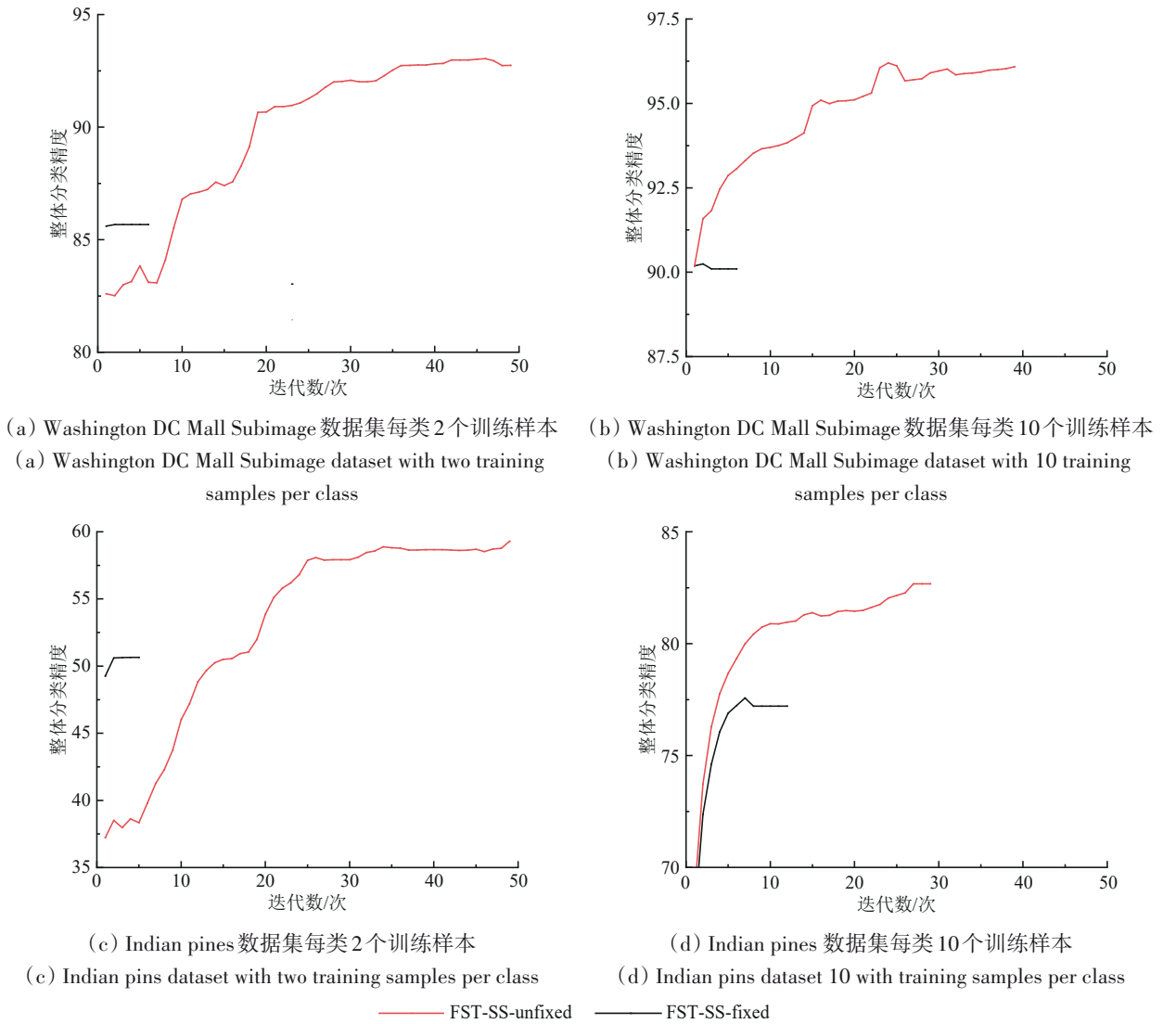


图 4 不同数据集下最大迭代次数参数分析图

Fig. 4 Parameter analysis diagram of the maximum number of iterations under different datasets

表 3 Washington DC Mall Subimage 数据集分类精度

Table 3 Classification accuracy of Washington DC Mall Subimage dataset

类别	1NN	Tri-Training	Star-SVM	ST-DP	LeMA	FST-SS-fixed	FST-SS-unfixed
C1	95.62	100	97.32	99.23	81.73	88.31	96.30
C2	73.46	59.56	93.29	70.66	92.50	94.23	88.27
C3	81.31	96.06	85.54	3.47	51.92	84.98	83.38
C4	93.04	83.59	97.04	44.86	16.56	92.89	95.69
C5	32.44	96.64	55.59	65.38	54.20	25.52	97.43
C6	45.47	15.44	58.21	66.18	82.84	97.92	99.75
C7	61.33	70.11	76.60	23.49	32.82	67.73	96.89
C8	96.99	0	96.80	91.25	97.65	98.40	95.39
OA	77.87	70.89	88.40	65.71	69.36	85.67	93.17
kappa	0.7357	0.6524	0.8589	0.5931	0.6258	0.8269	0.9179

注：表格中粗体为最好的结果；OA 代表整体分类精度；实验结果在每类选择 2 个训练样本下进行。

表4 Washington DC Mall Subimage数据集分类精度

Table 4 Classification accuracy of Washington DC Mall Subimage dataset

类别	1NN	Tri-Training	Star-SVM	ST-DP	LeMA	FST-SS-fixed	FST-SS-unfixed
C1	91.74	98.21	99.39	99.26	81.84	97.73	97.15
C2	88.10	95.72	94.03	94.44	92.31	89.19	96.45
C3	81.46	79.47	55.16	49.29	82.02	85.81	89.78
C4	92.08	95.15	95.78	84.93	95.36	92.54	93.39
C5	87.94	86.04	80.86	38.38	88.43	84.15	93.82
C6	67.33	69.80	72.15	89.48	84.53	74.01	97.28
C7	81.22	77.99	89.59	19.89	86.83	85.17	93.19
C8	96.02	95.92	96.68	91.47	96.49	93.93	98.20
OA	87.81	91.47	90.22	80.67	88.21	90.09	95.43
kappa	0.8534	0.8965	0.8813	0.7647	0.8655	0.8804	0.9446

注：粗体为最好结果；OA代表整体分类精度；实验结果在每类选择10个训练样本下进行。

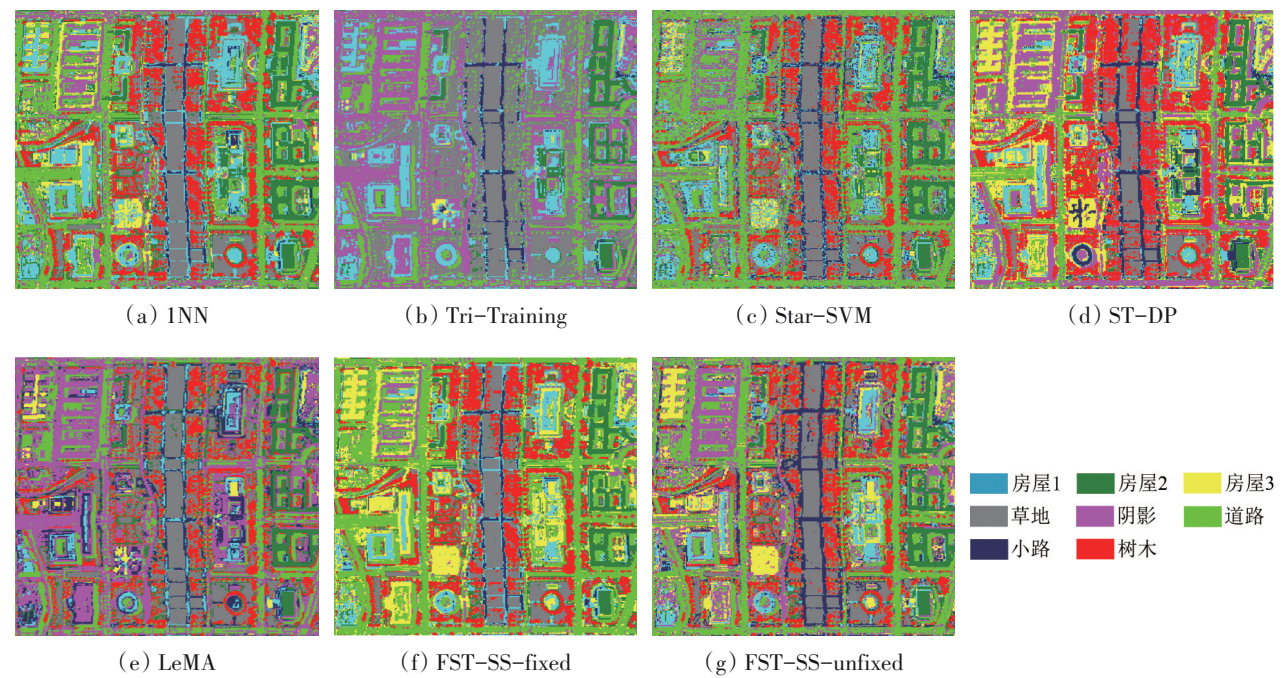
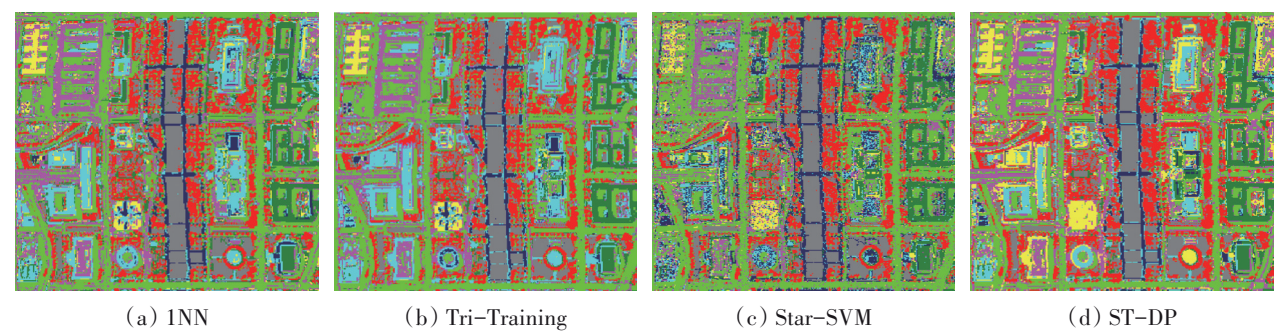


图5 Washington DC Mall Subimage数据集每类2个训练样本分类结果图

Fig.5 Classification results with the Washington DC Mall Subimage dataset with two training samples per class



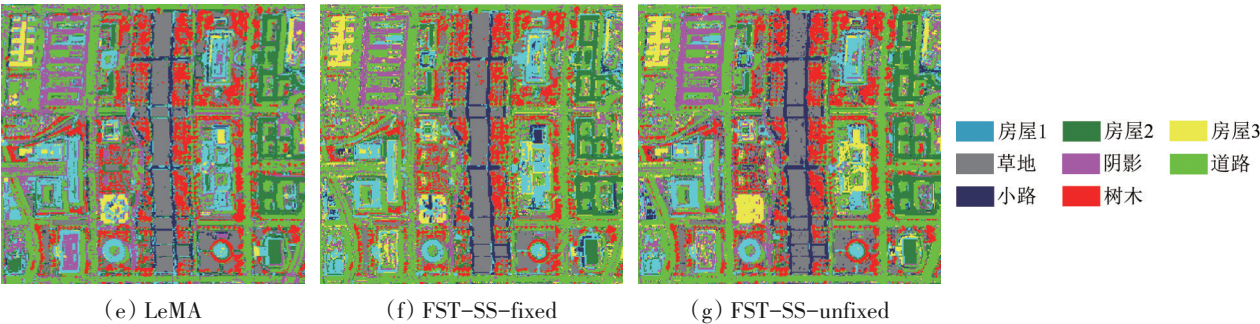


图6 Washington DC Mall Subimage数据集每类10个训练样本分类结果图
Fig. 6 Classification results with the Washington DC Mall Subimage dataset with 10 training samples per class

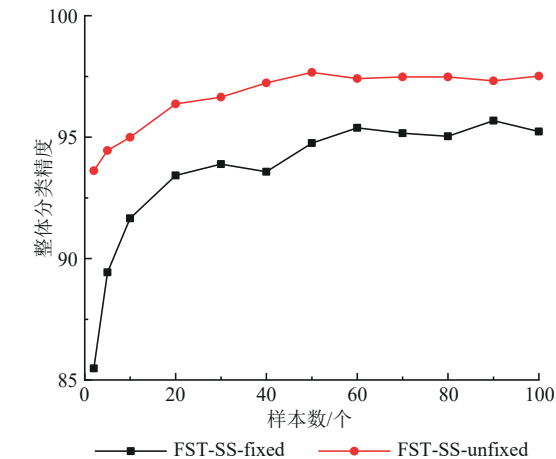


图7 Washington DC Mall Subimage数据集中不同标记样本下的整体分类精度OA图

Fig.7 Classification OA for different labeled samples in the Washington DC Mall Subimage dataset

3.3.3 计算效率分析

表7展示了不同算法的计算效率对比。从表7中可以发现FST-SS-fixed运行效率要略高于FST-SS-unfixed，这是因为FST-SS-fixed在训练过程中阈值的衰减速度要快于FST-SS-unfixed，所以FST-SS-fixed在训练过程中考虑标注的样本少于FST-SS-unfixed，因此FST-SS-fixed所需的计算时间更少，但是FST-SS-unfixed最终的计算精度要优于FST-SS-fixed。整体而言FST-SS-fixed和FST-SS-unfixed所需的计算时间都远远小于4种对比的半监督分类算法。这样验证了所提出算法的有效性。

表5 Indian Pines数据集分类精度
Table 5 Classification accuracy of Indian Pines dataset

								/%
类别	1NN	Tri-Training	Star-SVM	ST-DP	LeMA	FST-SS-fixed	FST-SS-unfixed	
C1	95.45	79.55	86.36	97.73	68.18	75.00	100	
C2	18.30	23.70	31.21	25.74	23.70	49.30	75.04	
C3	10.75	4.47	19.81	15.22	63.65	49.88	37.44	
C4	32.77	32.77	24.26	33.62	25.96	53.19	93.19	
C5	53.85	30.15	3.74	43.24	57.38	61.95	64.45	
C6	55.08	38.32	72.66	70.05	50.96	79.95	74.04	
C7	96.15	88.46	96.15	15.38	100	92.31	100	
C8	18.49	28.57	12.18	5.88	83.40	65.34	94.96	
C9	77.78	55.56	88.89	77.78	83.33	88.89	100	
C10	27.53	31.03	20.10	7.42	25.67	35.98	69.69	
C11	39.46	68.00	61.19	69.10	37.02	32.16	31.92	
C12	23.35	7.95	5.58	12.86	31.13	19.29	8.80	
C13	89.16	93.10	92.12	0.49	82.76	93.10	98.52	
C14	27.95	37.29	48.77	22.49	79.57	88.76	75.14	
C15	18.23	23.18	10.42	37.50	53.91	12.50	95.31	
C16	86.81	93.41	86.81	82.42	93.41	75.82	100	
OA	32.41	38.46	39.16	36.46	47.44	50.73	59.75	
kappa	0.2474	0.2999	0.3044	0.2717	0.4492	0.4492	0.5545	

注：粗体为最好的结果；OA代表整体分类精度；实验结果在每类选择2个训练样本下进行。

表 6 Indian Pines 数据集分类精度
Table 6 Classification accuracy of Indian Pines dataset

类别	INN	Tri-Training	Star-SVM	ST-DP	LeMA	FST-SS-fixed	FST-SS-unfixed
C1	75.00	77.78	72.22	72.22	97.22	97.22	100
C2	31.38	49.29	42.31	24.75	57.19	77.43	89.00
C3	38.90	40.24	46.95	32.56	54.02	68.90	80.98
C4	43.61	6.17	37.89	42.73	83.26	92.95	99.12
C5	76.53	72.09	80.34	81.40	94.71	90.70	95.77
C6	68.33	82.22	72.64	57.22	92.92	72.22	88.47
C7	100	94.44	100	88.89	83.33	0	0
C8	67.52	48.08	75.85	65.81	91.67	98.50	99.57
C9	90.00	100	50.00	50.00	100	100	100
C10	38.15	34.51	34.93	29.52	64.86	83.06	88.36
C11	51.74	44.38	57.22	51.41	46.58	66.99	69.24
C12	43.74	22.98	53.52	29.85	83.36	74.10	68.44
C13	89.74	87.69	96.41	86.15	99.49	96.41	98.97
C14	73.63	71.24	72.19	65.42	82.79	71.39	83.59
C15	27.33	48.67	57.71	22.34	79.52	74.73	98.67
C16	90.36	80.72	85.54	95.18	97.59	100	100
OA	52.05	50.76	57.55	46.92	68.50	75.78	83.16
kappa	0.4615	0.4445	0.5230	0.4015	0.6474	0.7252	0.8096

注：粗体为最好的结果；OA 代表整体分类精度；实验结果在每类选择 10 个训练样本下进行。

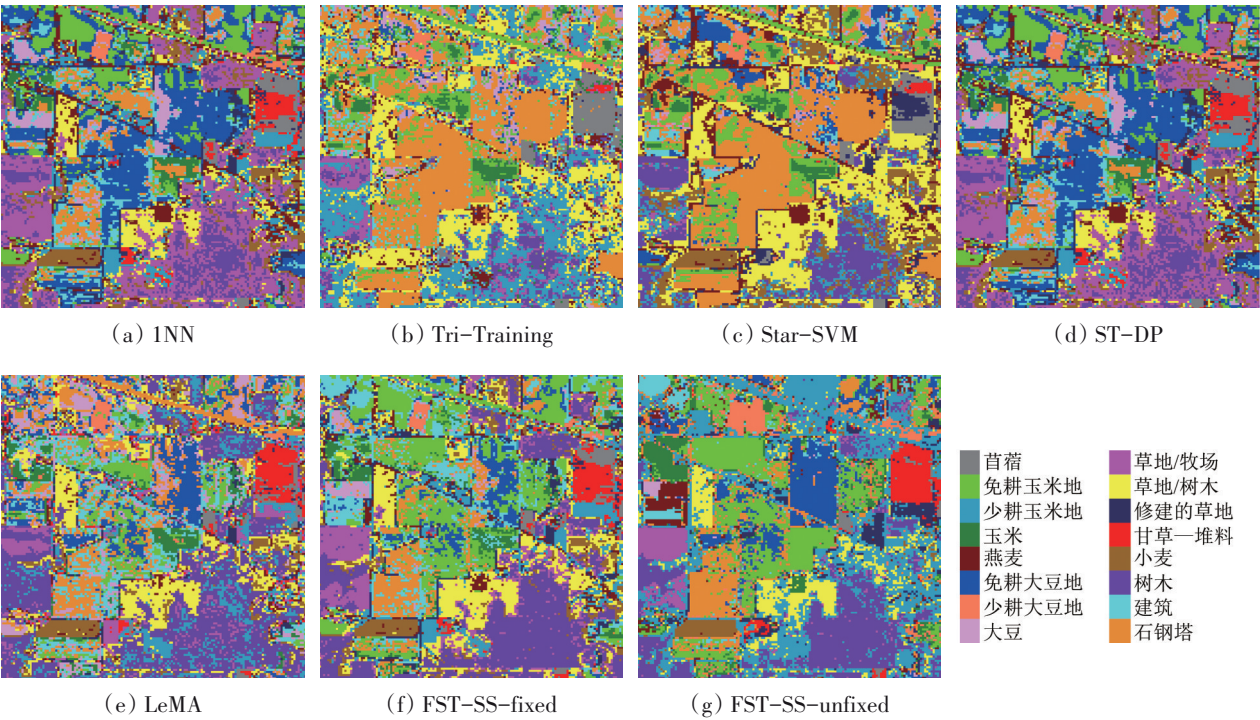


图 8 Indian Pine 数据集每类 2 个训练样本分类结果图

Fig.8 Classification results with the Indian Pine dataset with two training samples per class

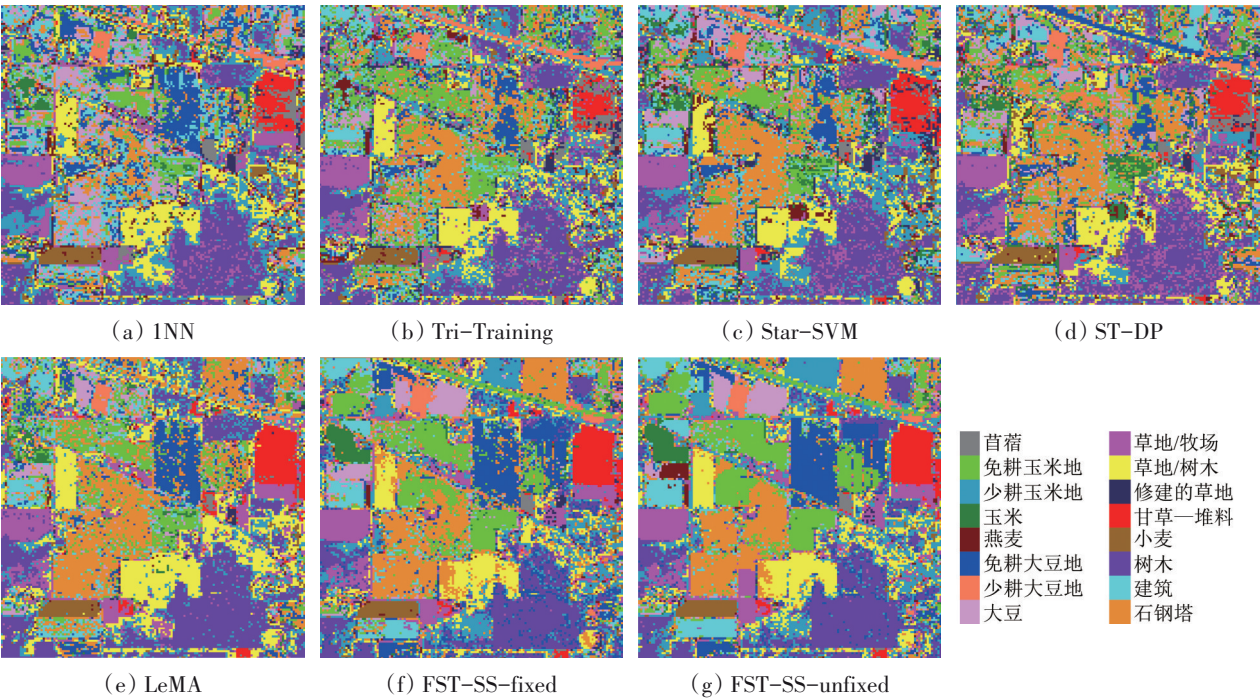


图9 Indian Pine数据集每类10个训练样本分类结果图
Fig. 9 Classification results with the Indian Pine dataset with ten training samples per class

表7 运行时间分析
Table 7 Run time analysis

数据集	每类训练样本数	运行时间/s						
		1NN	Tri-Training	Star-SVM	ST-DP	LeMA	FST-SS-fixed	FST-SS-unfixed
Washington DC Mall	2	0.22	562.6	54.8	1008.9	136.9	5.69	13.14
Subimage	10	0.39	156.9	199.5	1689.5	261.6	11.53	20.48
Indian Pines	2	0.27	612.25	45.56	1493.67	168.4	7.53	21.59
	10	0.45	184.48	169.83	2011.55	267.1	12.15	36.35

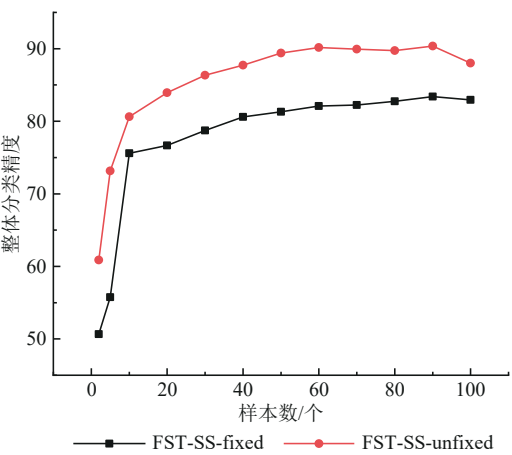


图10 Indian Pines数据集中不同标记样本下的整体分类精度
Fig.10 Classification OA for different labeled samples in the Indian Pines dataset

4 结 论

本文提出FST-SS算法来克服传统自训练方法

中标注精度低、训练效率低的问题。针对标注精度低的问题，FST-SS通过利用高光谱图像中地物空间分布的一致性，将高光谱图像中的空间信息作为补充，弥补了标记样本数目过少所导致的先验信息的不足，从而提高了对未标记样本的标注精度。而针对自训练方法训练效率低的问题，FST-SS没有沿用传统自训练学习使用分类器对未标记数据进行标记，而是利用空间—光谱信息对未标记样本进行多次筛选，然后对筛选出的样本直接赋予标记，避免了多次分类产生的高额时间成本，大大的提高了算法的计算效率。同时考虑到阈值在迭代过程中衰减速度过快，本文还设计了一种新的阈值计算方法减缓阈值的衰减。在两组真实的数据集上的实验结果证明，FST-SS相对于传统自训练方法在小样本条件下能取得更为优秀的分类结果，同时计算效率也远远优于传统自训练方法，证明了所提出算法的有效性。

此外, 本文提出的 FST-SS 方法还有进一步研究的空间。本文中自适应阈值的计算直接在原始欧式空间中进行, 而欧氏空间不具备良好的判别能力, 导致样本标注不够准确, 后续工作尝试将测度学习引入迭代训练过程中, 将数据从原始欧式空间映射到具有高判别力的测度空间后, 再进行样本标注, 通过测度学习提高标记数据的准确率。

参考文献 (References)

- Aydiv P S S and Minz S. 2018. Classification of hyperspectral images using self-training and a pseudo validation set. *Remote Sensing Letters*, 9(11): 1109-1117 [DOI: 10.1080/2150704x.2018.1511932]
- Bioucas-Dias J M, Plaza A, Camps-Valls G, Scheunders P, Nasrabadi N and Chanussot J. 2013. Hyperspectral remote sensing data analysis and future challenges. *IEEE Geoscience and Remote Sensing Magazine*, 1(2): 6-36 [DOI: 10.1109/mgrs.2013.2244672]
- Blaschke T. 2010. Object based image analysis for remote sensing. *ISPRS Journal of Photogrammetry and Remote Sensing*, 65(1): 2-16 [DOI: 10.1016/j.isprsjprs.2009.06.004]
- Bruzzzone L, Chi M and Marconcini M. 2006. A novel transductive SVM for semisupervised classification of remote-sensing images. *IEEE Transactions on Geoscience and Remote Sensing*, 44(11): 3363-3373 [DOI: 10.1109/tgrs.2006.877950]
- Camps-Valls G, Bandos Maracheva T V and Zhou D Y. 2007. Semi-supervised graph-based hyperspectral image classification. *IEEE Transactions on Geoscience and Remote Sensing*, 45(10): 3044-3054 [DOI: 10.1109/tgrs.2007.895416]
- Camps-Valls G, Gomez-Chova L, Munoz-Mari J, Vila-Frances J and Calpe-Maravilla J. 2006. Composite kernels for hyperspectral image classification. *IEEE Geoscience and Remote Sensing Letters*, 3(1): 93-97 [DOI: 10.1109/lgrs.2005.857031]
- Chen P H, Jiao L C, Liu F, Zhao J Q, Zhao Z Q and Liu S. 2017. Semi-supervised double sparse graphs based discriminant analysis for dimensionality reduction. *Pattern Recognition*, 61: 361-378 [DOI: 10.1016/j.patcog.2016.08.010]
- Cheung E and Li Y Y. 2017. Self-training with adaptive regularization for SVM//2017 International Joint Conference on Neural Networks (IJCNN). Anchorage: IEEE: 3633-3640 [DOI: 10.1109/IJCNN.2017.7966313]
- Dong Y N, Jin Y and Cheng S B. 2022. Clustered multiple manifold metric learning for hyperspectral image dimensionality reduction and classification. *IEEE Transactions on Geoscience and Remote Sensing*, 60: 5516813 [DOI: 10.1109/tgrs.2021.3123651]
- Gan H T, Li Z H, Wu W, Luo Z Z and Huang R. 2018. Safety-aware graph-based semi-supervised learning. *Expert Systems with Applications*, 107: 243-254 [DOI: 10.1016/j.eswa.2018.04.031]
- Ge H M, Pan H Z, Wang L G, Li C, Liu Y Z, Zhu W L and Teng Y P. 2021. A semi-supervised learning method for hyperspectral imagery based on self-training and local-based affinity propagation. *International Journal of Remote Sensing*, 42(17): 6391-6416 [DOI: 10.1080/01431161.2021.1934595]
- Ghamisi P, Plaza J, Chen Y S, Li J and Plaza A J. 2017. Advanced spectral classifiers for hyperspectral images: a review. *IEEE Geoscience and Remote Sensing Magazine*, 5(1): 8-32 [DOI: 10.1109/mgrs.2016.2616418]
- Gu S K and Jin Y C. 2017. Multi-train: a semi-supervised heterogeneous ensemble classifier. *Neurocomputing*, 249: 202-211 [DOI: 10.1016/j.neucom.2017.03.063]
- Gu X W. 2020. A self-training hierarchical prototype-based approach for semi-supervised classification. *Information Sciences*, 535: 204-224 [DOI: 10.1016/j.ins.2020.05.018]
- Hong D F, Yokoya N, Ge N, Chanussot J and Zhu X X. 2019. Learnable manifold alignment (LeMA): a semi-supervised cross-modality learning framework for land cover and land use classification. *ISPRS Journal of Photogrammetry and Remote Sensing*, 147: 193-205 [DOI: 10.1016/j.isprsjprs.2018.10.006]
- Jin Y, Dong Y N, Zhang Y X and Hu X Y. 2022. SSMD: dimensionality reduction and classification of hyperspectral images based on spatial-spectral manifold distance metric learning. *IEEE Transactions on Geoscience and Remote Sensing*, 60: 5538916 [DOI: 10.1109/tgrs.2022.3205178]
- Li F, Clausi D A, Xu L L and Wong A. 2018. ST-IRGS: a region-based self-training algorithm applied to hyperspectral image classification and segmentation. *IEEE Transactions on Geoscience and Remote Sensing*, 56(1): 3-16 [DOI: 10.1109/tgrs.2017.2713123]
- Li Z Y, Togo R, Ogawa T and Haseyama M. 2020. Chronic gastritis classification using gastric X-ray images with a semi-supervised learning method based on tri-training. *Medical and Biological Engineering and Computing*, 58(6): 1239-1250 [DOI: 10.1007/s11517-020-02159-z]
- Lu X C, Zhang J P, Li T and Zhang Y. 2016. A novel synergetic classification approach for hyperspectral and panchromatic images based on self-learning. *IEEE Transactions on Geoscience and Remote Sensing*, 54(8): 4917-4928 [DOI: 10.1109/tgrs.2016.2553047]
- Plaza A, Benediktsson J A, Boardman J W, Brazile J, Bruzzzone L, Camps-Valls G, Chanussot J, Fauvel M, Gamba P, Gualtieri A, Marconcini M, Tilton J C and Trianni G. 2009. Recent advances in techniques for hyperspectral image processing. *Remote Sensing of Environment*, 113: S110-S122 [DOI: 10.1016/j.rse.2007.07.028]
- Tan K, Zhu J S, Du Q, Wu L X and Du P J. 2016. A novel tri-training technique for semi-supervised classification of hyperspectral images based on diversity measurement. *Remote Sensing*, 8(9): 749 [DOI: 10.3390/rs8090749]
- van Engelen J E and Hoos H H. 2020. A survey on semi-supervised learning. *Machine Learning*, 109(2): 373-440 [DOI: 10.1007/s10994-019-05855-6]
- Wang F and Zhang C S. 2008. Label propagation through linear Neighborhoods. *IEEE Transactions on Knowledge and Data Engineering*

- ing, 20(1): 55-67 [DOI: 10.1109/tkde.2007.190672]
- Wang L G, Hao S Y, Wang Q M and Wang Y. 2014. Semi-supervised classification for hyperspectral imagery based on spatial-spectral Label Propagation. *ISPRS Journal of Photogrammetry and Remote Sensing*, 97: 123-137 [DOI: 10.1016/j.isprsjprs.2014.08.016]
- Wu D, Shang M S, Luo X, Xu J, Yan H Y, Deng W H and Wang G Y. 2018. Self-training semi-supervised classification based on density peaks of data. *Neurocomputing*, 275: 180-191 [DOI: 10.1016/j.neucom.2017.05.072]
- Zerguine A, Shafi A and Bettayeb M. 2001. Multilayer perceptron-based DFE with lattice structure. *IEEE Transactions on Neural Networks*, 12(3): 532-545 [DOI: 10.1109/72.925556]
- Zhang L, Zhou W D and Jiao L C. 2004. Wavelet support vector machine. *IEEE Transactions on Systems, Man, and Cybernetics, Part B (Cybernetics)*, 34(1): 34-39 [DOI: 10.1109/tsmcb.2003.811113]
- Zhou Z H and Li M. 2005. Tri-training: exploiting unlabeled data using three classifiers. *IEEE Transactions on Knowledge and Data Engineering*, 17(11): 1529-1541 [DOI: 10.1109/tkde.2005.186]
- Zoidi O, Tefas A, Nikolaidis N and Pitas I. 2018. Positive and negative label propagations. *IEEE Transactions on Circuits and Systems for Video Technology*, 28(2): 342-355 [DOI: 10.1109/tcsvt.2016.2598671]

Fast self-training based on spatial-spectral information for hyperspectral image classification

JIN Yao^{1,2,3}, DONG Yanni^{1,2,4}, DU Bo⁵

1. Institute of Geophysics and Geomatics, China University of Geosciences, Wuhan 430074, China;

2. Hubei Luojia Laboratory, Wuhan 430079, China;

3. State Key Laboratory of Information Engineering in Surveying, Mapping and Remote Sensing, Wuhan University, Wuhan 430079, China;

4. School of Resource and Environmental Science, Wuhan University, Wuhan 430079, China;

5. School of Computer Science, Wuhan University, Wuhan 430072, China

Abstract: Hyperspectral image classification has been a popular issue in the field of hyperspectral image interpretation. The prominent problem in hyperspectral image classification currently involves the considerably time-consuming and expensive manual acquisition of labeled samples for hyperspectral images in practical applications. This problem leads to a sparse number of training samples and increases the difficulty of obtaining good classification results. Self-training methods are widely used in hyperspectral image classification to solve the difficulty in labeled sample acquisition in hyperspectral image classification. Traditional self-training methods mostly use spectral information to classify unlabeled samples and then utilize the expanded labeled data set to iteratively train the classifier to complete the classification task. In this model, the spatial information provided by the hyperspectral images is ignored, resulting in poor classification accuracy. Simultaneously, the classification of unlabeled data must be completed once during each iteration, resulting in a significant time cost. Therefore, a fast self-training method based on spatial - spectral information is proposed in this paper for hyperspectral image classification to address the above problems.

FST-SS (Fast Self-training based on Spatial-Spectral information) supplements the spatial information in hyperspectral images by exploiting the consistency of the spatial distribution of features in the hyperspectral images. Instead of using the classifier to classify unlabeled samples, this approach uses the spatial - spectral information to filter unlabeled data and extend the labeled samples during the iterative process. The spatial nearest neighbors are first selected using a spatial domain patch for the initial labeled sample. The spatial nearest neighbors are then filtered using an adaptive threshold to obtain the spatial - spectral nearest neighbors to be labeled. Finally, the classifier is trained to complete the classification task based on the expanded labeled samples.

This paper compares FST-SS with the supervised classification algorithm 1NN, the semi-supervised classification algorithm Star-SVM, Tri-Training, ST-DP, and LeMA on two real hyperspectral datasets to demonstrate the effectiveness of FST-SS. Experimental results show that the overall classification accuracy reaches 93.17% and 95.43% when 2 and 10 training samples, respectively, are selected for each class in the Washington DC Mall subimage dataset. The overall classification accuracy in Indian Pines dataset reaches 59.75% and 86.13%, which is a significant improvement compared with the comparison algorithm.

The FST-SS algorithm uses the spatial - spectral information provided by hyperspectral images to label unlabeled samples by combining the ideas of self-training methods. Compared with the conventional self-training methods, instead of using a classifier for classification, FST-SS uses the spatial - spectral information to filter the unlabeled samples directly, which markedly improves the computational efficiency of the algorithm.

Key words: hyperspectral image classification, semi-supervised classification, spatial-spectral information, self-training

Supported by National Natural Science Foundation of China (No. 62222116, 62171417, 41871243, 62141112); Project Supported by the Open Fund of Hubei Luojia Laboratory (No. 220100058)