



(12) 发明专利

(10) 授权公告号 CN 108197643 B

(45) 授权公告日 2021. 11. 30

(21) 申请号 201711447557.X

(56) 对比文件

(22) 申请日 2017.12.27

CN 105787513 A, 2016.07.20

US 2014029839 A1, 2014.01.30

(65) 同一申请的已公布的文献号

申请公布号 CN 108197643 A

审查员 王菲

(43) 申请公布日 2018.06.22

(73) 专利权人 佛山科学技术学院

地址 528000 广东省佛山市南海区狮山镇

仙溪水库西路佛山科学技术学院

(72) 发明人 易长安 顾艳春 王东 李晓东

胡小生 何志敏

(74) 专利代理机构 广东广信君达律师事务所

44329

代理人 杨晓松

(51) Int. Cl.

G06K 9/62 (2006.01)

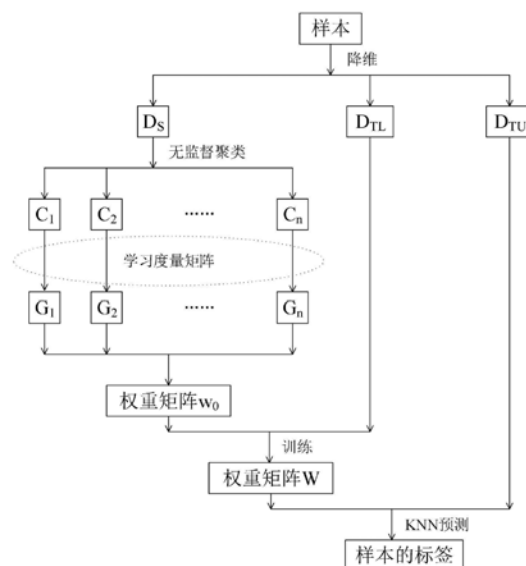
权利要求书2页 说明书6页 附图3页

(54) 发明名称

一种基于无监督聚类 and 度量学习的迁移学习方法

(57) 摘要

本发明涉及一种基于无监督聚类 and 度量学习的迁移学习方法, 先准备源域样本 D_S 和目标域样本 D_T ; 再使用主成份分析方法降低全部样本的维度; 然后对源域样本 D_S 进行无监督聚类, 将源域样本聚为多个类别; 再为每个聚类学习一个度量矩阵 G , 通过每个聚类和所有的度量矩阵的关联性得出权重矩阵 w_0 , 并根据已知标签的目标域样本 D_{TL} 训练权重矩阵 w_0 , 得到最优的权重矩阵 W , 最后根据权重矩阵 W 预测未知标签的目标域样本 D_{TU} 的标签。本发明使用主成份分析方法降低样本的维度, 降低后续计算的复杂度。使用无监督聚类方法将样本聚为多个类别, 更能够反应样本的本质特性。使用带标签的目标域样本来学习权重矩阵, 更符合目标域样本的实际情况。



1. 一种基于无监督聚类 and 度量学习的迁移学习方法, 其特征在于, 包括以下步骤:

S1、准备源域样本 D_S 和目标域样本 D_T ;

S2、使用主成份分析方法降低全部样本的维度;

S3、对源域样本 D_S 进行无监督聚类, 将源域样本聚为多个类别, 本质相似的样本聚在一起;

S4、为每个聚类学习一个度量矩阵 G ;

S5、通过每个聚类和所有的度量矩阵的关联性得出权重矩阵 w_0 ;

S6、根据已知标签的目标域样本 D_{TL} 训练权重矩阵 w_0 , 得到最优的权重矩阵 W ;

S7、根据权重矩阵 W 预测未知标签的目标域样本 D_{TU} 的标签;

所述步骤S6的具体过程如下:

S61、求出每个已知标签的目标域样本 D_{TL} 的所属聚类 C_i ;

S62、使用每个度量矩阵进行预测, 若结果与真实标签相同, 则不需要调整, 若结果与真实标签不相同, 则调整聚类 C_i 和该度量矩阵所对应元素的权重; 直至所有已知标签的目标域样本 D_{TL} 均已测试;

S63、若经所有已知标签的目标域样本 D_{TL} 测试后的权重矩阵 w_0 满足收敛条件, 则训练结束, 该权重矩阵即为最优的权重矩阵 W ; 否则继续使用 D_{TL} 更新权重矩阵, 直至满足收敛条件为止。

2. 根据权利要求1所述的一种基于无监督聚类 and 度量学习的迁移学习方法, 其特征在于: 所述步骤S4每个聚类在学习一个度量矩阵之前均需要将样本的顺序打乱。

3. 根据权利要求1所述的一种基于无监督聚类 and 度量学习的迁移学习方法, 其特征在于: 所述步骤S4学习度量矩阵 G 之前需设置收敛条件:

$||G - G_0|| < \text{Lamda}$, 其中, G_0 为度量矩阵 G 的初始值, 将初始值设置为对称单位矩阵, Lamda 为设定的阈值, $||G - G_0||$ 表示 G 和 G_0 之间距离;

求解度量矩阵 G 的目标函数的公式如下:

$$\min_G \int p(x; G_0) \log \frac{p(x; G_0)}{p(x; G)} dx$$

subject to $d_G(x_i, x_j) \leq \alpha, x_i, x_j \in S,$

$d_G(x_i, x_j) \geq \beta, x_i, x_j \in D$

其中, x_i, x_j 指其中的某个样本, S 表示 x_i 与 x_j 同类, D 表示 x_i 与 x_j 不同类; $p(x; G)$ 表示度量矩阵 G 下的样本 x 的概率分布; $p(x; G_0)$ 表示度量矩阵 G_0 下的样本 x 的概率分布;

$d_G(x_i, x_j) = \sqrt{(x_i - x_j)^T G (x_i - x_j)}$; α 和 β 为设定的阈值。

4. 根据权利要求3所述的一种基于无监督聚类 and 度量学习的迁移学习方法, 其特征在于: 所述阈值 α 和 β 的设置方式为: 从训练集里任意选取100个样本对, 并且将样本两两之间的距离按从小到大的顺序排列, 排第5的距离值为 α , 排第95的距离值为 β 。

5. 根据权利要求1所述的一种基于无监督聚类 and 度量学习的迁移学习方法, 其特征在于: 所述权重矩阵 w_0 满足收敛条件后, 对于 w_0 的每个元素 w_{ij} , $1 \leq i, j \leq n$, n 为聚类数目, 均执行归一化操作。

6.根据权利要求1所述的一种基于无监督聚类 and 度量学习的迁移学习方法,其特征在于:所述步骤S7的具体过程如下:

S71、求出每个未知标签的目标域样本 D_{tu} 各自所属的聚类 C_t ;

S72、使用全部的度量矩阵预测样本i的类型;

S73、使用KNN分类器进行分类得出分类值;

S74、若分类值大于阈值threshold,则预测样本i的标签值为Label1,否则预测样本i的标签值为Label2,Label1大于Label2;其中,

$threshold = (Label1 - Label2) / NumOfCategory$, NumOfCategory为标签种类数目。

一种基于无监督聚类 and 度量学习的迁移学习方法

技术领域

[0001] 本发明涉及迁移学习的技术领域,尤其涉及到一种基于无监督聚类 and 度量学习的迁移学习方法。

背景技术

[0002] 迁移学习用于解决训练样本(源域)和测试样本(目标域)的分布不一致问题。知识迁移在现实生活中广泛存在,例如,学会打羽毛球之后要学习打网球,学会削苹果之后要学习削梨子,学会辨识猫之后要学习如何辨识狗,这些知识相似、但不相同。如何最大化地利用已有的知识,从而以最快、最有效的方式来学习新知识,是迁移学习的主要难点。人类需要有知识迁移的能力,机器人也是如此。RoboBrain是由来自于斯坦福大学、康奈尔大学等高校的教授主导的项目,其核心目标就是为机器人赋予知识迁移的能力。人类之所以能够不断掌握新知识,其主要原因就是能够举一反三、迁移知识,这种能力也是“机器换人、智能制造”要解决的难题之一。深度学习的优势是从复杂环境里提取特征,但是却不具备迁移能力。AlphaGo的核心是深度学习和强化学习,但是,如果比赛前突然改变围棋的盘面大小,AlphaGo就无法适应,而人类却可以自如地应对。近年来,迁移学习已经成为人工智能、机器人等领域的核心研究问题。

[0003] 已有的迁移学习技术主要是基于特征、基于实例、或者基于度量。基于特征的方法尝试通过学习特征或子空间来匹配源域和目标域。基于实例的方法主要是调整源域实例的权重从而减少源域和目标域的数据分布之间的差异。这些技术以欧氏距离为基础,不能够获得样本维度之间的相关性。但是,样本的不同维度之间却存在关联,例如,人的身高和体重、人的性别和爱好之间可能存在某种关联,这种关联会影响到整体的特征。度量学习考虑了样本的维度之间的相关性。基于度量学习的迁移学习方法相对较少,已有的方法仅仅考虑样本的标签,忽视了样本特征所蕴含的本质属性。例如,动物的体重可以划分为多个档次,如small、medium、big等等。30~50公斤的猪(样本)可以算作small,但是,同样体重的狗(样本)却可以算作medium。所以,在这种情况下样本的特征更能够反应其本质属性。因此,将聚类与度量学习相结合,更有利于在迁移情况下预测样本的标签。

发明内容

[0004] 本发明的目的在于克服现有技术的不足,提供一种基于无监督聚类 and 度量学习的迁移学习方法,其既考虑样本维度之间的相关性、又考虑样本特征所隐含的本质属性,使源域知识能更好地迁移到目标域。

[0005] 为实现上述目的,本发明所提供的技术方案为:

[0006] 包括以下步骤:

[0007] S1、准备源域样本 D_S 和目标域样本 D_T ;

[0008] 源域样本指过去的样本,目标域样本指当前要预测的样本。由于这两类样本的分布存在一定的区别,所以不可以仅仅使用源域样本来训练预测模型。另外,如果对全部的目

标域样本 D_T 都人工贴标签,其工作量太大。所以,本方案是先给一小部分的目标域样本贴标签,然后结合 D_S 来为 D_T 训练新的预测模型。

[0009] 所述目标域样本 D_T 分成两部分: $D_T = \{D_{TL} \cup D_{TU}\}$,就样本数量来说, $D_{TL} \ll D_{TU}$ 。 D_{TL} 的标签已知,该标签由人工赋予, D_{TU} 的标签未知。训练集为 T , $T = \{D_S \cup D_{TL}\}$ 。

[0010] S2、使用主成份分析方法降低全部样本的维度;

[0011] 在图像识别等领域,样本往往高达数万维甚至更多,计算度量矩阵也就变得十分复杂。如何既有效地降低样本的维度、又保留原始数据的本质特征,这是该步骤需要解决的问题。

[0012] S3、对源域样本 D_S 进行无监督聚类,将源域样本聚为多个类别,本质相似的样本聚在一起;

[0013] S4、为每个聚类学习一个度量矩阵 G ;对于每个聚类,在学习度量矩阵之前都需要将样本的顺序打乱,其目的是让选中的样本更具有随机性。

[0014] S5、通过每个聚类和所有的度量矩阵的关联性得出权重矩阵 w_0 ;

[0015] S6、根据已知标签的目标域样本 D_{TL} 训练权重矩阵 w_0 ,得到最优的权重矩阵 W ;

[0016] S7、根据权重矩阵 W 预测未知标签的目标域样本 D_{TU} 的标签。

[0017] 进一步地,步骤S4中,度量矩阵 G 为对称矩阵,其初始值 G_0 设置为单位矩阵,对角线元素为1,其余全部为0。 G_0 在经过许多次迭代更新后就趋于稳定。学习度量矩阵 G 之前需设置收敛条件: $\|G - G_0\| < \text{Lamda}$,其中, G_0 为度量矩阵 G 的初始值, Lamda 为设定的阈值, $\|G - G_0\|$ 表示 G 和 G_0 之间距离;

[0018] 求解度量矩阵 G 的目标函数的公式如下:

$$[0019] \quad \min_G \int p(x; G_0) \log \frac{p(x; G_0)}{p(x; G)} dx$$

[0020] subject to $d_G(x_i, x_j) \leq \alpha (i, j) \in S,$

[0021] $d_G(x_i, x_j) \geq \beta (i, j) \in D$

[0022] 该公式反应了:不同类的样本之间距离大、同类样本之间的距离小。

[0023] 其中, x_i, x_j 指其中的某个样本, S 表示 x_i 与 x_j 同类, D 表示 x_i 与 x_j 不同类; $p(x; G)$ 表示度量矩阵 G 下的样本 x 的概率分布; $p(x; G_0)$ 表示度量矩阵 G_0 下的样本 x 的概率分布;

$d_G(x_i, x_j) = \sqrt{(x_i - x_j)^T G (x_i - x_j)}$; α 和 β 为设定的阈值。

[0024] 阈值 α 和 β 的设置方式为:从训练集里任意选取100个样本对,并且将它们之间的距离按从小到大的顺序排列,排第5的距离值为 α ,排第95的距离值为 β 。

[0025] 进一步地,步骤S6的具体过程如下:

[0026] S61、求出每个已知标签的目标域样本 D_{TL} 的所属聚类 C_i ;

[0027] S62、使用每个度量矩阵进行预测,若结果与真实标签相同,则不需要调整,若结果与真实标签不相同,则调整聚类 C_i 和该度量矩阵所对应元素的权重;直至所有已知标签的目标域样本 D_{TL} 均已测试;

[0028] S63、若经所有已知标签的目标域样本 D_{TL} 测试后的权重矩阵 w_0 满足收敛条件,则训练结束,该权重矩阵即为最优的权重矩阵 W ;否则继续使用 D_{TL} 更新权重矩阵,直至满足收敛条件为止。

[0029] 权重矩阵 w_0 满足收敛条件后,对于 w_0 的每个元素 w_{ij} ($1 \leq i, j \leq n$) 均执行归一化操作,公式如下:

$$[0030] \quad w_{ij} = \frac{w_{ij}}{\sum_{k=1}^n w_{ik}} \quad \circ$$

[0031] 进一步地,步骤S7的具体过程如下:

[0032] S71、求出每个未知标签的目标域样本 D_{TU} 各自所属的聚类 C_t ;

[0033] S72、使用全部的度量矩阵预测样本 i 的类型;

[0034] S73、使用KNN分类器进行分类得出分类值;

[0035] S74、若分类值大于阈值threshold,则预测样本 i 的标签值为Label1,否则预测样本 i 的标签值为Label2;其中,

[0036] $\text{threshold} = (\text{Label1} - \text{Label2}) / \text{NumOfCategory}$, NumOfCategory为标签种类数目。

[0037] 与现有技术相比,本方案原理如下:

[0038] 先准备源域样本 D_S 和目标域样本 D_T ,再使用主成份分析方法降低全部样本的维度,然后对源域样本 D_S 进行无监督聚类,将源域样本聚为多个类别,本质相似的样本聚在一起,再为每个聚类学习一个度量矩阵 G ,通过每个聚类 and 所有的度量矩阵的关联性得出权重矩阵 w_0 ,并根据已知标签的目标域样本 D_{TL} 训练权重矩阵 w_0 ,得到最优的权重矩阵 W ,最后根据权重矩阵 W 预测未知标签的目标域样本 D_{TU} 的标签。

[0039] 与现有技术相比,本方案优点如下:

[0040] 1、使用主成份分析方法来降低样本的维度,从而降低后续计算的复杂度。

[0041] 2、使用无监督聚类方法将样本聚为多个类别,从而更能够反应样本的本质特性。

[0042] 3、使用带标签的目标域样本来学习权重矩阵,从而更符合目标域样本的实际情况。

附图说明

[0043] 图1为本发明实施例的流程图;

[0044] 图2为本发明实施例中权重矩阵的训练流程图;

[0045] 图3为本发明实施例中预测目标域样本标签的流程图。

具体实施方式

[0046] 下面结合具体实施例对本发明作进一步说明:

[0047] 参见附图1所示,本实施例所述的一种基于无监督聚类和度量学习的迁移学习方法,包括以下步骤:

[0048] S1、准备源域样本 D_S 和目标域样本 D_T ;

[0049] 目标域样本 D_T 分成两部分: $D_T = \{D_{TL} \cup D_{TU}\}$, $D_{TL} \ll D_{TU}$, D_{TL} 的标签已知,该标签由人工赋予, D_{TU} 的标签未知。训练集为 T , $T = \{D_S \cup D_{TL}\}$;

[0050] S2、使用主成份分析方法降低全部样本的维度;

[0051] S3、对源域样本 D_S 进行无监督聚类,将源域样本聚为多个类别,本质相似的样本聚

在一起;

[0052] S4、每个聚类学习一个度量矩阵G;对于每个聚类,在学习度量矩阵之前都需要将样本的顺序打乱,如使用randperm函数(以Matlab为例),其目的是让选中的样本更具有随机性;另外,设置收敛条件: $\|G - G_0\| < \text{Lamda}$,其中, G_0 为度量矩阵G的初始值,Lamda为设定的阈值, $\|G - G_0\|$ 表示G和 G_0 之间距离;

[0053] 求解度量矩阵G的目标函数的公式如下:

$$[0054] \quad \min_G \int p(x; G_0) \log \frac{p(x; G_0)}{p(x; G)} dx$$

[0055] subject to $d_G(x_i, x_j) \leq \alpha (i, j) \in S,$

[0056] $d_G(x_i, x_j) \geq \beta (i, j) \in D$

[0057] 其中, x_i, x_j 指其中的某个样本,S表示 x_i 与 x_j 同类,D表示 x_i 与 x_j 不同类; $p(x; G)$ 表示度量矩阵G下的样本x的概率分布; $p(x; G_0)$ 表示度量矩阵 G_0 下的样本x的概率分布;

$d_G(x_i, x_j) = \sqrt{(x_i - x_j)^T G (x_i - x_j)}$; α 和 β 为设定的阈值。

[0058] 阈值 α 和 β 的设置方式为:从训练集里任意选取100个样本对,并且将它们之间的距离按从小到大的顺序排列,排第5的距离值为 α ,排第95的距离值为 β 。

[0059] S5、通过每个聚类和所有的度量矩阵的关联性得出权重矩阵 w_0 ;每个聚类都和所有的度量矩阵相关联,如表1所示:

[0060]

聚类 \ 度量矩阵	G_1	G_2	$\dots$	G_n
C_1	w_{11}	w_{12}	$\dots$	w_{1n}
C_2	w_{21}	w_{22}	$\dots$	w_{2n}
$\dots$	$\dots$	$\dots$	$\dots$	$\dots$
C_n	w_{n1}	w_{n2}	$\dots$	w_{nn}

[0061] 表1

[0062] 权重矩阵 $w_0 = \{w_{ij}, i = 1, 2, \dots, n, j = 1, 2, \dots, n\}$,如表2所示:

[0063]

w_{11}	w_{12}	$\dots$	w_{1n}
w_{21}	w_{22}	$\dots$	w_{2n}

[0064]

$\dots$	$\dots$	$\dots$	$\dots$
w_{n1}	w_{n2}	$\dots$	w_{nn}

[0065] 表2

[0066] n为聚类数目；

[0067] 权重矩阵 w_0 的初始值为对角矩阵,如表3所示:

[0068]

1	0	...	0
0	1	...	0
...	...	...	...
0	0	0	1

[0069] 表3

[0070] S6、根据已知标签的目标域样本 D_{TL} 训练权重矩阵 w_0 ,得到最优的权重矩阵 W ;如图2所示,具体步骤如下:

[0071] S61、求出每个已知标签的目标域样本 D_{TL} 的所属聚类 C_i ;

[0072] S62、使用每个度量矩阵进行预测,若结果与真实标签相同,则不需要调整,若结果与真实标签不相同,则调整聚类 C_i 和该度量矩阵所对应元素的权重;直至所有已知标签的目标域样本 D_{TL} 均已测试;

[0073] S63、若经所有已知标签的目标域样本 D_{TL} 测试后的权重矩阵 w_0 满足收敛条件,则训练结束,该权重矩阵即为最优的权重矩阵 W ;否则继续使用 D_{TL} 更新权重矩阵,直至满足收敛条件为止。

[0074] 其中, w_0 的收敛条件为 $||w_{s+1}-w_s|| < \text{Sigma}$,Sigma为一个很小的正数, w_s 表示第s次迭代后的权重矩阵;

[0075] 权重矩阵 w_0 满足收敛条件后,对于 w_0 的每个元素 w_{ij} ($1 \leq i, j \leq n$),均执行归一化操作,公式如下:

[0076]

$$W_{ij} = \frac{W_{ij}}{\sum_{k=1}^n W_{ik}};$$

[0077] S7、根据权重矩阵 W 预测未知标签的目标域样本 D_{TU} 的标签;如图3所示(假设样本的标签值为0、1),具体步骤如下:

[0078] S71、求出每个未知标签的目标域样本 D_{TU} 各自所属的聚类 C_t ;

[0079] S72、使用全部的度量矩阵分别预测样本i的类型为 m_j ($1 \leq j \leq n$);

[0080] S73、使用KNN分类器进行分类得出分类值:

[0081] $w_i = w_{t1}m_1 + w_{t2}m_2 + \dots + w_{tn}m_n$;

[0082] S74、若分类值大于阈值threshold,则预测样本i的标签值为Label1,否则预测样本i的标签值为Label2 (Label1大于Label2);其中,threshold = (Label1-Label2) / NumOfCategory,NumOfCategory为标签种类数目。

[0083] 如Label1和Label2分别为1和0,NumOfCategory为2,可得threshold = (1-0) / 2 = 0.5;若分类值大于阈值0.5,则预测样本i的标签值为1,否则预测样本i的标签值为0。

[0084] 本实施例可以用于常见的迁移学习数据集,例如:

[0085] i) 20Newsgroup (文本分类, <http://people.csail.mit.edu/jrennie/20Newsgroups/>)

[0086] ii) Office-Caltech (物体识别, <http://www.eecs.berkeley.edu/~jhoffman/domainadapt/>)

[0087] iii) USPS (手写字体识别, <http://www-i6.informatik.rwth-aachen.de/~keyzers/usps.html>) 和 MNIST (手写字体识别, <http://yann.lecun.com/exdb/mnist>)

[0088] 其中, USPS、MNIST 的数据分布不同但相关, 两者均可以作为源域或目标域数据。

[0089] 而且本实施例具有以下优点:

[0090] 1、使用主成份分析方法降低样本的维度, 从而降低后续计算的复杂度。

[0091] 2、使用无监督聚类方法将样本聚为多个类别, 从而更能够反应样本的本质特性。

[0092] 3、使用带标签的目标域样本来学习权重矩阵, 从而更符合目标域样本的实际情况。

[0093] 以上所述之实施例子只为本发明之较佳实施例, 并非以此限制本发明的实施范围, 故凡依本发明之形状、原理所作的变化, 均应涵盖在本发明的保护范围内。

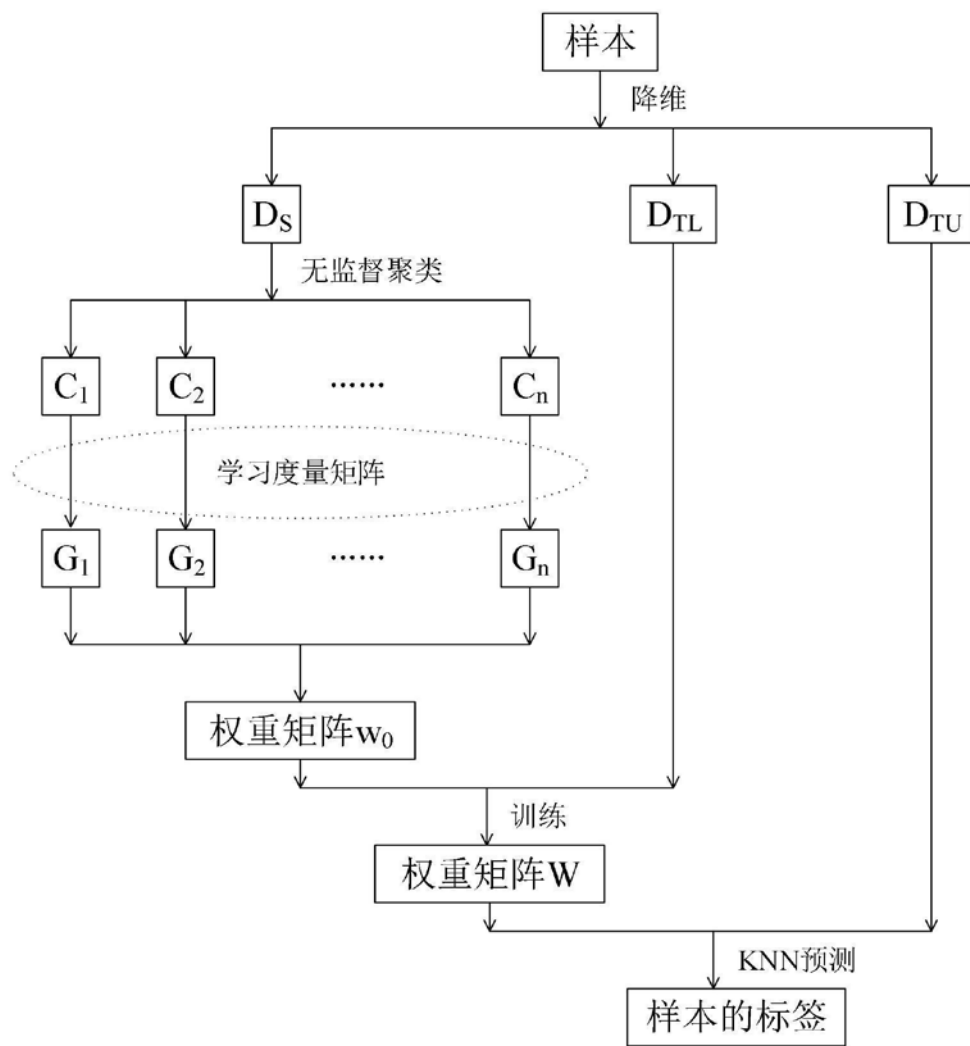


图1

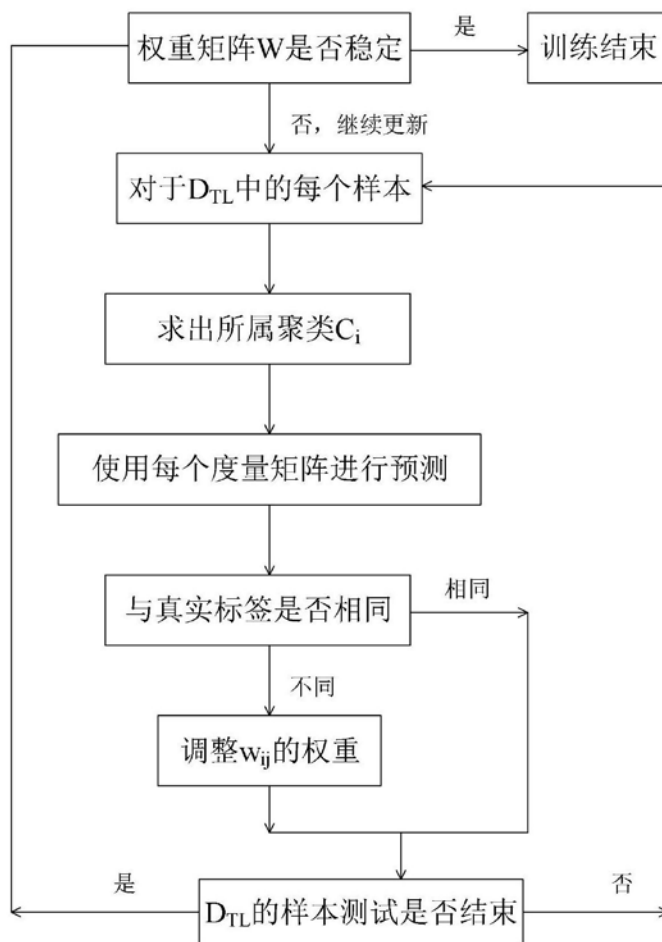


图2

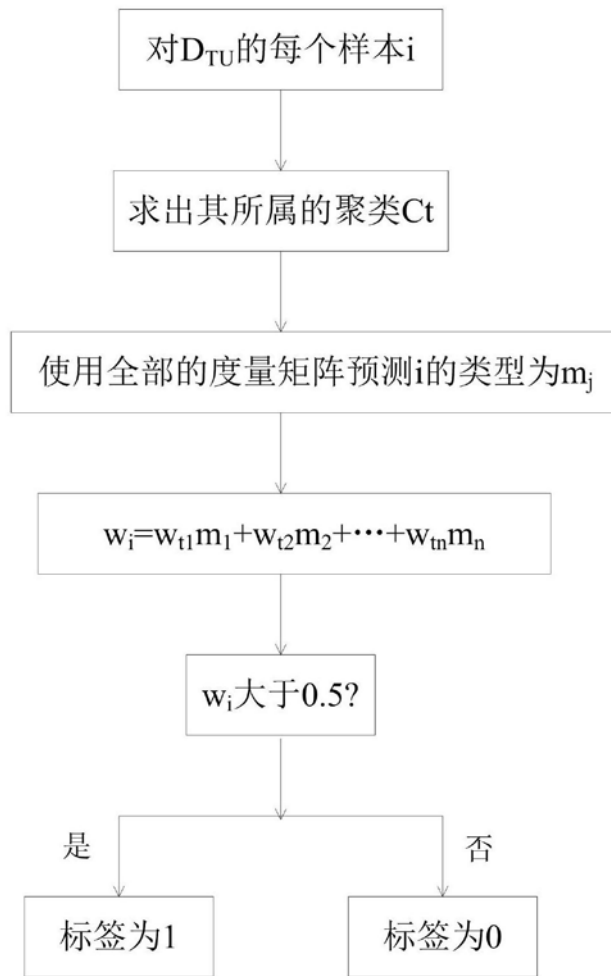


图3