

基于不同损失和距离函数的乘更新分类算法*

胡晓雄, 贾育秦[†]

(太原科技大学 机械工程学院, 太原 030024)

摘要: 与批处理的学习算法相比较,在线预测的算法由于在大样本集上具有预测准确性高、时间的空间复杂度小,因此占有非常大的优势。对于 Warmuth 与 Jivinen 提出的保守性和权衡正确性的在线学习框架,已经得到相当广泛的应用,但是在他们所提出的指数梯度下降算法以及梯度下降算法中,对于目标函数中的损失函数,在其求导过程中使用近似步骤将会造成在线学习结果的恶化。根据对偶的最优化理论,基于不同损失函数与距离函数,提出了变更新分类算法。通过一系列的实验分析与研究,结果表明,该算法使得预测准确率得到了很大的提高,从而验证了该算法的正确性、可行性和有效性。

关键词: 乘权更新; 在线学习; 非近似更新; 分类算法

中图分类号: TP393; TP181

文献标志码: A

文章编号: 1001-3695(2014)02-0344-04

doi:10.3969/j.issn.1001-3695.2014.02.005

Multiplier update classification algorithm based on different loss and distance function

HU Xiao-xiong, JIA Yu-qin[†]

(School of Mechanical Engineering, Taiyuan University of Science & Technology, Taiyuan 030024, China)

Abstract: In comparison with batch learning algorithm, online prediction algorithms on large sample sets with high predictive accuracy, low time and space complexity, and therefore, that plays a very big advantage. The conservative and weigh the correctness online learning framework proposed by Warmuth and Jivinen, had gotten a wide range of applications. But in the Warmuth & Jivinen's index gradient down algorithms and gradient descent algorithm, for the loss function in the objective function, the approximate steps would result in the deterioration of the results of the online learning during the derivation process. Based on the dual optimization theory, this paper proposed changes to the new classification algorithm based on a variety of loss function and distance function. Through a series of experimental analysis and research results show that the algorithm makes the prediction accuracy better, and verifies the correctness, feasibility and effectiveness of the algorithm.

Key words: multiply weight updating; online learning; non-approximately update; classification algorithm

0 引言

预测与训练是离线算法的两个阶段,当处于训练阶段时,所有的训练样本被算法一次输入,训练获得判别函数;当处于预测阶段时,所有的训练样本再一次被输入,并实施预测。具体来讲就是,倘若有样本集 $(x_i, y_i)_{i=1}^m, 1 \leq i \leq m$, 训练的样本集被算法输入,并实施训练获得 $f(\cdot)$ (判别函数),而后使用 $f(\cdot)$ 输入预测的样本集并实施预测。与离线算法相比较,在线算法在训练大量样本时,其时间空间的复杂度小^[1-12],同时对于现实生活中的不少问题,其全是随机训练样本,从而依次达到在线序列学习,一次性获得所有训练的样本集是不可能的。在在线算法学习的过程中,每次输入一个样本,算法依据判别准则对样本标签进行预测,同时获得样本的真实类标签,而后算法依据损失函数对损失进行计算。当预测有错误时,依据权值的修正准则对权值实施修正;当预测正确时,对权值不进行修正。通过在线算法的权值更新过程可知,由于在在线算法学习的过程中,其每次都输入一个样本,时间空间的复杂度就显得相对较小,因此可以较好地应用于对大样本集的预测。

并且在预测的过程中对权值进行修正,即当预测有错误时,对判别准则不断地进行修正,从而在下次预测时达到不再预测错误的目的,最终可以使得判别准则实现最优化。具体地,对于在线算法,它是按轮进行的,倘若有样本集 $(x_i, y_i), 1 \leq i \leq m$, 当在第 i 轮时,算法获得一个样本 $x_i \in \mathbb{R}^d$,算法依据预测准则 $y = \text{sign}(w^T x)$ 对样本类标签 $\hat{y} \in \{-1, +1\}$ 进行预测,同时算法返回样本的损失函数 $L(y_i, \hat{y}_i)$ 与真实类标签 $y_i \in \{-1, +1\}$,损失函数可以选取为平方损失 $L = (y - \hat{y})^2$ 与铰链损失 $L = \max\{0, 1 - y_i(w^T x_i)\}$ 。当预测错误时 $y_i \neq \hat{y}_i$,对权值 $w_{i+1} = w_i + a_i y_i x_i$ (其中 a_i 为非负系数)进行修正,直至预测结束,获得最终的判别准则函数以及权值 w ;当预测正确时 $y_i = \hat{y}_i$,对权值不进行修正,即 $w_{i+1} = w_i$ 。

Kivinen 等人^[13]提出了一种新的在线回归算法,他们把目标函数定义为

$$C(w) = d(\cdot) + \eta L(\cdot) \quad (1)$$

其中: $d(\cdot)$ 为 Bregman 距离;

$$d_F(u, w') = d_F(u, w) + d_F(u, w') + (f(w') - f(w)) \cdot (w - u)$$

收稿日期: 2013-05-03; **修回日期:** 2013-06-19 **基金项目:** 国家自然科学基金资助项目(61273127, 61074077); 国家“973”计划资助项目(2012CB312901, 2012CB312905); 国家“863”计划资助项目(2011AA01A103)

作者简介: 胡晓雄(1988-),男,山西交城人,硕士研究生,主要研究方向为 CIMS; 贾育秦(1958-),男(通信作者),山西太原人,教授,硕士,主要研究方向为 CIMS(tyhmiqq@163.com)。

$L(\cdot)$ 是损失函数, 对函数 $C(w) = d(\cdot) + \eta L(\cdot)$ 进行最小化, 获得权值的更新, 即 $w_{i+1} = \arg \min C(w)$, 于是使得预测正确率以及权值更新的近似度有了保证。倘若 $C(w)$ 为凸函数, 并且其可导, 令 $\nabla C(w) = 0$, 则 $w_{i+1} = w_i - \eta \nabla C(w_i)$ 。

Warmuth 和 Kivinen 提出使用 $C(w)$ 在 w_i 处的梯度值, 近似地替代 $C(w)$ 在 w_{i+1} 处的梯度值, 则令 $\nabla C(w_{i+1}) = \nabla C(w_i)$, 可得 $w_{i+1} = w_i - \eta \nabla C(w_i)$ 。

若选用 $d(w_k, w_{k-1}) = \sum_{i=1}^n (w_{k-1,i} - w_{k-1,i})^2$, $L(y_k, y_k) = (y_k - w_k^T x_k)^2$, 获得加权的更新公式为 $w_{k+1} = w_k - 2\eta(y_{k-1} - y_k)x_k$ 。

若假设 $w_k, w_{k-1} \in [0, \infty)^n$ 是概率向量, 即 $\sum_{i=1}^n w_{k,i} = 1$, 选用相关熵的距离 $d_{re}(w_k, w_{k-1}) = \sum_{i=1}^n w_{k,i} \ln \frac{w_{k,i}}{w_{k-1,i}}$; 其中进一步规定, 当 $w_{i,j} = 0$ 时, $0 \ln 0 = 0$, 当 $w_{k-1,i} = 0$ 时, $d_{re}(w_k, w_{k-1}) = \infty$, 此时对式(1)进行求解, 即可获得乘权更新形式为

$$w_{k,i} = \frac{r_{k-1,i} w_{k-1,i}}{\sum_{p=1}^n r_{k-1,p} w_{k-1,p}}$$

其中: $r_{k-1,p} = e^{-m(\hat{y}_{k-1} - y_{k-1})x_{k-1,p}}$ 。

Warmuth 和 Kivinen 在对目标函数求导的过程中, 采用近似步骤造成的问题作了细致的阐释。在 $z_k = w_k^T x_k$ 点用 $z_k = w_{k-1}^T x_k$ 代替对 $\frac{\partial L(y_k, z_k)}{\partial z_k} = L'(z_k)$ 进行求导, 当在采用相关熵的损失函数 $L(y, \hat{y}) = y \ln \left(\frac{y}{\hat{y}} \right) + (1-y) \ln \frac{1-y}{1-\hat{y}}$ 时, 就会造成非常严重的问题, 此时, $L'(y, z) = -\frac{y_k}{z_k} + \frac{1-y_k}{1-z_k}$ 。对于 $0 < y_k < 1$, 当 z_k 接近于 1 或者 0 时, $L'(y_k, z_k)$ 的值将会发生非常大的变化, 也就是说, 当预测值 $w_{k-1}^T x_k$ 接近于 1 或者 0 时, $L'(y_k, w_{k-1}^T x_k)$ 的值就可能会与真实值 $L'(y_k, w_k^T x_k)$ 相差非常大, 因此会导致算法求出的导数值相当不精确, 对算法的分类结果造成很大的影响。

虽然本文也采用了 Warmuth 和 Kivinen 损失函数以及距离函数的双目标优化模式, 但是本文与 Warmuth 和 Kivinen 的不同之处和改进有: a) 在部分权值的更新中, 给出了 η 的迭代选取规则, 而在 W&K 的算法中, 对于 η 的选取并没有给出选取的规则, 是人为地指定参数, 通过实际的实验分析与研究表明, 倘若选取得不正确, 对于分类算法的预测性能将会大大地降低; b) 在算法的迭代过程中, 对权值 w 进行求导时, 不采用近似更新的模式, 从而消除了文献[1]中所指出的问题, 同时引入了新的损失函数及距离函数, 因此获得了新的权值更新公式, 其算法是分类算法, 而不是回归算法。

1 乘权更新在线分类算法研究

在本文的研究过程中, 采用的距离和损失函数如下:

1) 距离函数

a) 归一化的相关熵的距离 $d(\cdot) = \sum_{i=1}^n w_i \ln \frac{w_i}{s_i}$

b) 平方距离 $d(\cdot) = \|w - s\|^2$

c) 未归一化的相关熵的距离

$$d(\cdot) = \sum_{i=1}^n s_i - w_i + w_i \ln \frac{w_i}{s_i}$$

2) 损失函数

a) 相关熵的损失

$$L(y, \hat{y}) = y \ln \left(\frac{y}{\hat{y}} \right) + (1-y) \ln \frac{1-y}{1-\hat{y}}$$

b) 平方损失 $L(\cdot) = (y - w^T x)^2$

c) 铰链损失 $L(\cdot) = \max\{0, 1 - y(w^T x)\}$

定理 1 当在目标函数 $C(w) = d(\cdot) + \eta L(\cdot)$ 中, 损失函数是铰链的损失函数 $L(\cdot) = \max\{0, 1 - y(w^T x)\}$, 距离函数是未归一化的相关熵的距离 $d(\cdot) = \sum_{i=1}^n s_i - w_i + w_i \ln \frac{w_i}{s_i}$ 时, 权值的更新公式为 $w = se^{\lambda xy}$, 其中: $\eta = \frac{-\ln xy}{xy}$ 。

证明 $C(w) = \sum_{i=1}^n s_i - w_i + w_i \ln \frac{w_i}{s_i} + \eta[1 - y(w^T x)]$

通过求导可得

$$\frac{\partial C(w)}{\partial w_i} = -1 + \ln \frac{w_i}{s_i} + 1 + \eta(-y_i x_i) = 0$$

$$w = se^{\eta xy}$$

把权值的更新准则代入到 $C(w)$:

$$C(w) = s - se^{\eta xy} + se^{\eta xy} \eta xy + \eta(1 - yse^{\eta xy} x) = s - se^{\eta xy} + \eta$$

对 η 进行求导且令其等于 0, 于是获得学习率的更新公式为

$$\frac{\partial C(w)}{\partial \eta} = 1 - se^{\eta xy} xy = 0, \eta = \frac{-\ln xy}{xy}$$

定理 2 当在目标函数 $C(w) = d(\cdot) + \eta L(\cdot)$ 中, 损失函数是铰链的损失函数 $L(\cdot) = \max\{0, 1 - y(w^T x)\}$, 距离函数是归一化的相关熵的距离函数 $d(\cdot) = \sum_{i=1}^n w_i \ln \frac{w_i}{s_i}$, 则权值的更新公式为 $w = se^{\eta xy - 1}$, 其中 $\eta = \frac{1 - \ln sxy}{xy}$ 。

证明 $C(w) = \sum_{i=1}^n w_i \ln \frac{w_i}{s_i} + \eta[1 - y(w^T x)]$

对 w_i 进行求导可得

$$\frac{\partial C(w)}{\partial w_i} = 1 + \sum_{i=1}^n \ln \frac{w_i}{s_i} + \eta(-x_i y_i) = 0$$

$$w = se^{\eta xy - 1}$$

把权值更新代入 $C(w)$ 中:

$$C(w) = \sum_{i=1}^n se^{\eta xy - 1} \ln e^{\eta xy - 1} + \eta[1 - y(se^{\eta xy - 1} x)] =$$

$$\sum_{i=1}^n se^{\eta xy - 1} (\eta xy - 1) + \eta(1 - sxy e^{\eta xy - 1}) = -se^{\eta xy - 1} + \eta$$

对 η 进行求导且令其等于 0, 于是获得学习率的更新公式为

$$\frac{\partial C(w)}{\partial \eta} = -se^{\eta xy - 1} (xy) + 1 = 0$$

$$\eta xy - 1 = -\ln sxy, \eta = \frac{1 - \ln sxy}{xy}$$

定理 3 当在目标函数 $C(w) = d(\cdot) + \eta L(\cdot)$ 中, 损失函数是相关熵的损失函数 $L(y, \hat{y}) = y \ln \left(\frac{y}{\hat{y}} \right) + (1-y) \ln \frac{1-y}{1-\hat{y}}$, 距离函数是归一化的距离函数 $d(\cdot) = \sum_{i=1}^n w_i \ln \frac{w_i}{s_i}$, 则当 $y = 0$ 时, 权值的更新公式为 $w = s \cdot \exp(-\lambda x - 1)$; 当 $y = 1$ 时, 权值的更新公式为 $w = s \cdot \exp(\lambda x - 1)$ 。

证明 $C(w) = \sum_{i=1}^n w_i \ln \frac{w_i}{s_i} + \eta \left(y \ln \left(\frac{y}{\hat{y}} \right) + (1-y) \ln \frac{1-y}{1-\hat{y}} \right)$

假设样本标签是 $y \in \{0, 1\}$, 分成以下两种情况进行讨论:

a) 当 $y = 0$ 时,

$$C(w) \big|_{y=0} = \sum_{i=1}^n w_i \ln \frac{w_i}{s_i} + \eta(-\ln(1-\hat{y}))$$

$$\frac{\partial C(w)}{\partial w_i} \big|_{y=0} = \ln \frac{w_i}{s_i} + 1 + \left(\frac{-\eta x}{1 - w^T x} \right)$$

因为 $C(w)$ 是凸函数, 并且存在 w , 使得 $[-\ln(1-w^T x)] < 0$, 所以存在 w^* 与 λ^* , $C(w)$ 取得最优值时, 可以使得互补松弛的条件 $\eta[-\ln(1-w^T x)] = 0$ 。

当 $C(w)$ 取得最优值时, λ 是 λ^* , $\lambda^*[-\ln(1-(w^*)^T x)] = 0$, $(1-(w^*)^T x) = 1$, 则 $\frac{\partial C(w)}{\partial w_i} \Big|_{y=0} = \ln \frac{w_i}{s_i} + 1 + \left(\frac{-\eta x}{1-w^T x} \right) = 0$ 变为 $\ln \frac{w_i}{s_i} + 1 + \eta x = 0$, 因此权值最终更新为 $w = s \cdot \exp(-\eta x - 1)$ 。

b) 当 $y = 1$ 时,

$$C(w) \Big|_{y=1} = \sum_{i=1}^n w_i \ln \frac{w_i}{s_i} + \eta(-\ln(w^T x))$$

$$\frac{\partial C(w)}{\partial w} \Big|_{y=1} = \ln \frac{w_i}{s_i} + 1 + \eta \left(\frac{-x_i}{w^T x} \right)$$

同理, 存在 w^* 与 η^* , 使得 $C(w)$ 最小时符合松弛条件为 $-\eta(-\ln(w^T x)) = 0$, $w^T x = 1$ 。获得 $\ln \frac{w_i}{s_i} + 1 - \eta x_i = 0$, 因此最终的权值更新为 $w = s \cdot \exp(\eta x - 1)$ 。

一般地, 一个最优化的数学模型可以表示为

$$\begin{aligned} \min f(x) \\ \text{s. t. } h_j(x) &= 0 \quad j=1, \dots, p \\ g_k(x) &\leq 0 \quad k=1, \dots, q \\ x &\in X \subset \mathbb{R}^n \end{aligned}$$

当且仅当对每一个 x^* 都符合 KKT 最优优化条件时取得最优解:

$$\text{a) } h_j(x^*) = 0, j=1, \dots, p; g_k(x^*) \leq 0, k=1, \dots, q.$$

$$\text{b) } \nabla f(x^*) + \sum_{j=1}^p \lambda_j \nabla h_j(x^*) + \sum_{k=1}^q \mu_k \nabla g_k(x^*) = 0, \lambda_j \neq 0, \mu_k \geq 0, \mu_k g_k(x^*) = 0.$$

通过以上分析可知, 对于原样本集以及违背 KKT 条件的新增样本集, 其在学习过程中是需要重点考虑的。同时为了尽可能地减少丢失原样本集中的有用信息, 还需要考虑符合 KKT 条件的样本中与原分类间隔距离比较近的样本。

2 算法的实现

上一章中提出了 RERE 算法(基于相关熵的损失函数与相关熵的距离函数的权值更新在线算法)、REH 算法(基于相关熵的距离函数与铰链损失函数的权值更新在线算法)、UREH 算法(基于未归一化的相关熵的距离函数与铰链损失函数的权值更新在线算法), 表 1 是变更新的分类算法的系数 λ 的选取以及权值更新公式。

表 1 算法的权值更新和系数 λ

变更新算法	权更新公式	λ
RERE	$y=0, w = s \cdot \exp(-\lambda x - 1),$ $y=1, w = s \cdot \exp(\lambda x - 1)$	—
REH	$w = se^{\lambda xy - 1}$	$\lambda = \frac{1 - \ln sxy}{xy}$
UREH	$w = se^{\lambda xy}$	$\lambda = \frac{-\ln xy}{xy}$

算法 1

初始化: $w_1 = s$ 。

输入: 样本集 $(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)$ 与 λ (使用 RERE 算法时)。

a) 当算法收到样本 x_i 后, 计算准则 $\hat{y}_i = \text{sign}(w_i^T x_i)$ 。

b) 误差的计算:

(a) $L_i = \max\{0, 1 - y_i(w_i^T x_i)\}$ (使用 REH 算法与 UREH 算法);

(b) $L_i = (y_i - \hat{y}_i)^2$ (使用 RERE 算法)。

当 $L_i = 0, w_{i+1} = w_i$ 。

当 $L_i \neq 0$, 更新权值为

(a) $w = se^{\lambda xy}$ (使用 RERE 算法), $\lambda = \frac{-\ln xy}{xy}$;

(b) $w = se^{\lambda xy - 1}$ (使用 REH 算法), $\lambda = \frac{1 - \ln sxy}{xy}$;

(c) $y=0, w = s \cdot \exp(-\lambda x - 1)$ (使用 RERE 算法); $y=1, w = s \cdot \exp(\lambda x - 1)$ 。

输出: 权值及预测误差。

3 实验分析与研究

通过采用 UCI (<http://archive.ics.uci.edu/me/datasets.html>) 数据中的 USPS、Letter 与 MNIST 数据实施预测, 以此来验证算法的预测的效果, 对于本文提出的三种预测算法, 将其与现存的 MU 算法进行对比与研究。表 2 所示是三种数据的信息。

表 2 三个实际数据集的特征

数据集	类别标签个数	特征个数
MNIST	10	780
Letter	26	16
USPS	10	780

在进行实验时, 首先使得数据线性归一化, 计算各个样本

x_i 的分量之和 $\sum_{j=1}^n x_{i,j}$, 并令 $x_i = \frac{x_i}{\sum_{j=1}^n x_{i,j}}$, 同时通过算法 RERE、REH

权值的更新公式可知, 对于线性归一化后的数据, 其在算法 RERE 及 REH 的权值的更新将会非常慢。本文引入了优化过程中的步长 η , 以此来解决权值更新太慢的问题, 则算法 REH 的权值更新为 $w = \eta se^{\lambda xy - 1}$, $w = \eta se^{-1} e^{\lambda xy}$ 。令 $\eta^* = \eta e^{-1}$, 则 $w = \eta^* se^{\lambda xy}$, 其中 η^* 是固定值。

同理, 可以得到 RERE 算法的权值更新为

$$y=0, w = \eta^* s \cdot \exp(-\lambda x)$$

$$y=1, w = \eta^* s \cdot \exp(\lambda x)$$

在实验进行的过程中, 对于分类预测的效果可以使用预测的正确率来表示。以下是算法进行的过程: 每次输入一个 x_i (训练样本), 通过使用权值的更新公式可以获得一个 w (权值), 同时根据预测的训练是否正确, 不断地更新权值, 直至样本训练结束, 获得判别准则 f 以及最终的 w (权值), 然后针对测试的样本使用判别准则 f 对其实施测试。

3.1 USPS 数据实验

本文选取 USPS 数据集中的四种数据(1 和 2 类、1 和 10 类、3 和 8 类、7 和 10 类), 图 1 所示是 USPS 数据的示例图。

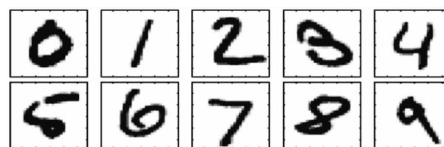


图 1 USPS 数据集

在本实验中, 选取 $\eta^* = 0.9$, λ 的值如表 3 所示, 对四种数据的训练如图 2 所示。由对图 2(a)(c)(d) 的分析可知, RERE 与 REH 是重合的, MU 的训练误差是最差的, 与 UREH 相比, RERE 和 REH 的训练误差比较差, UREH 的训练误差是最好的; 由图 2(b) 可知, RERE 与 REH 是重合的, MU 的训练误差是最差的, 与 RERE 和 REH 相比, UREH 的训练误差比较差。通过对表 4 分析可知, 使用 RERE 和 REH 对四种数据的预测正确率是一致的, 当在相同的条件下, MU 的预测正确率是最低的, RERE 与 REH 的预测正确率是最高的, UREH 的预测正确率比 MU 要高, 但是比 RERE 和 REH 要低。

表3 USPS数据集 λ 的值

	1与2类	1与10类	3与8类	7与10类
λ	10.648 801	15.229 730	12.292 356	14.442 339

表4 USPS数据集测试准确率

算法	1与2类	1与10类	3与8类	7与10类
MU	0.576 2	0.669 8	0.669 8	0.589 9
RERE	0.855 5	0.990 7	0.964 1	0.953 5
REH	0.855 5	0.990 7	0.964 1	0.953 5
UREH	0.825 0	0.957 1	0.959 4	0.951 0

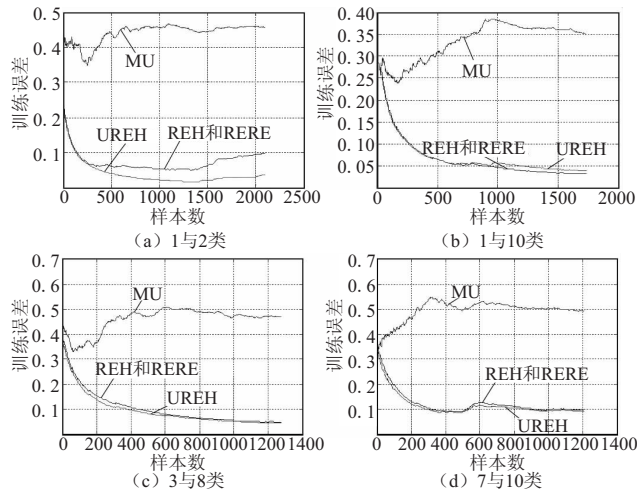


图2 USPS训练误差

3.2 MNIST数据实验

本文选取 MNIST 数据集中的四种数据(1 与 2 类、1 与 10 类、3 与 8 类、7 与 10 类)。MNIST 数据集是美国国家标准和技术研究所提供的 10 000 个测试样本与 60 000 个训练样本的手写字符的数据库, $20 \times 20 = 400$ 像素空间中的向量。图 3 所示是 MNIST 数据集。在 MNIST 数据实验的过程中,选取 $\eta^* = 0.9$, λ 的值如表 5 所示,由图 4(a) ~ (d) 可知, RERE 与 REH 是重合的, MU 的训练误差最大, UREH 的训练误差最小,与 UREH 相比, RERE 与 REH 的训练误差是比较差的。

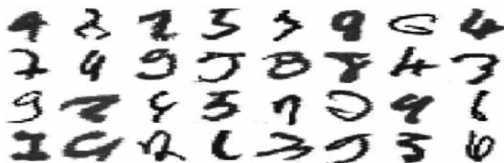


图3 MNIST数据集

表5 MNIST数据集 λ 的值

	1与2类	1与10类	3与8类	7与10类
λ	6.959 772	6.959 772	10.919 305	12.029 907

通过对表 6 的分析可知,在对四种数据预测的过程中, MU 最终的预测正确率是最低的,对于 RERE、REH 及 UREH 的最终预测正确率基本上都是相等的。

表6 MNIST数据集测试准确率

算法	1与2类	1与10类	3与8类	7与10类
MU	0.523 8	0.536 6	0.509 1	0.512 0
RERE	0.976 9	0.943 7	0.959 1	0.977 0
REH	0.976 9	0.943 7	0.959 1	0.977 0
UREH	0.978 3	0.940 8	0.953 1	0.982 5

3.3 Letter实验

本文选取 Letter 数据集中的四种数据(1 与 2 类、1 与 22 类、7 与 17 类、17 与 18 类),在 Letter 的数据实验过程中,选取

$\eta^* = 0.9$, λ 的值如表 7 所示,对四种数据的训练如图 5 所示。通过对图 5 分析可知, RERE 与 REH 是重合的,与 UREH 相比, RERE 与 REH 的训练误差比较差, UREH 的训练误差最好, MU 的训练误差最差,并且比本文所提出的三种算法要差很多。

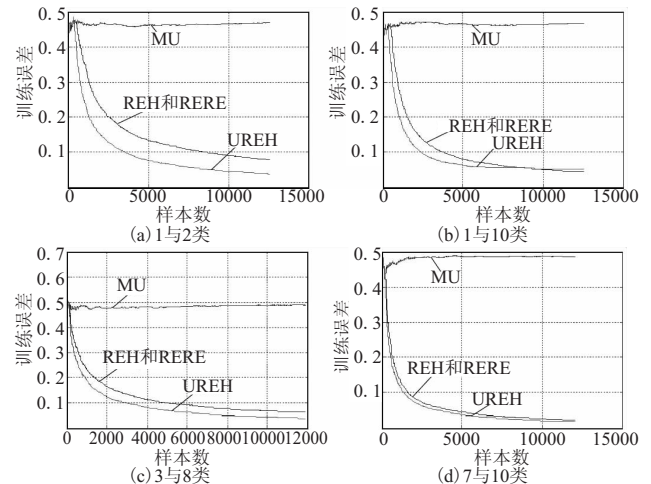


图4 MNIST数据集训练误差

表7 Letter数据集 λ 的值

	1与2类	1与22类	7与17类	17与18类
λ	2.029 221	2.286 585	2.001 838	2.180 769

通过对表 8 分析可知,当在相同的条件下时, MU 的预测效果是最差的,并且 MU 的预测误差在 50% 左右,说明 MU 算法起不到有效的分类作用, RERE 与 REH 的预测误差是一致的,并且预测的效果最好, UREH 预测的效果比 MU 要好,但是比 RERE 和 REH 差。

表8 Letter数据集测试准确率

算法	1与2类	1与22类	7与17类	17与18类
MU	0.503 2	0.484 7	0.474 6	0.500 0
RERE	0.968 6	0.990 8	0.834 9	0.970 2
REH	0.968 6	0.990 8	0.834 9	0.970 2
UREH	0.968 2	0.825 2	0.623 1	0.967 3

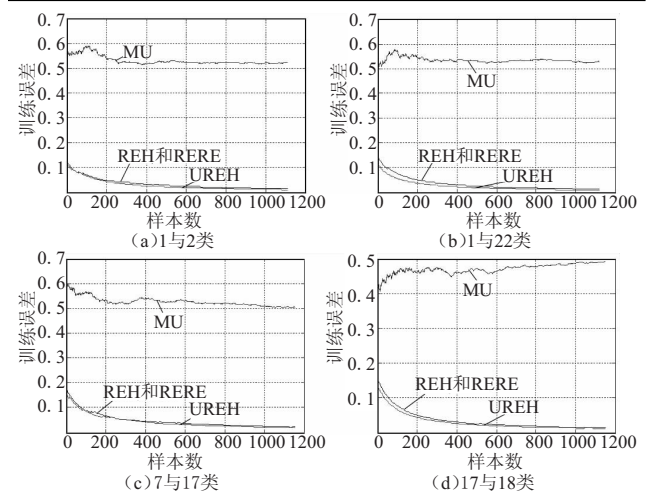


图5 Letter训练误差

4 结束语

基于 Warmuth 和 Kivinen 对偶的优化理论所提出的权值更新模式,本文使用不同的距离函数及损失函数提出了三种在线的分类算法,同时使用各种 UCI 数据集实施了分类实验研究。通过实验表明,所提出的权值更新模式具有正确性、可行性与有效性。下一步的主要任务是把正则化引入到对权值的更新目标函数中,针对高维稀疏数据对其实施在线的分类处理。

(下转第 383 页)

销量先接近然后差距又不断增大。 GC 契约下的系统期望效用和集中决策下的期望效用均是随着 λ 的增加而减小,两者的变化趋势与最优供货量的变化趋势是一样的。

当制造商的风险厌恶系数与零售商的风险厌恶系数不同时,制造商的风险厌恶水平和零售商的风险厌恶水平对供需系统的影响是不同的。当制造商的风险厌恶水平不变时,随着零售商的风险厌恶水平的增加,最优供货量、努力水平、批发价格以及系统的期望效用都是减小的;而当零售商的风险厌恶水平不变时,随着制造商的风险厌恶水平的增加,批发价格和制造商的期望效用是增加的,其余决策变量是减小的。

当供需系统的风险厌恶水平高于制造商和零售商的风险厌恶水平时, GC 契约下的供需系统可以达到甚至超过集中决策下的供需系统的期望效用;而当系统的风险厌恶水平不变时,制造商和零售商的风险厌恶系数越高, GC 契约对协调供需系统的作用越小,且发现零售商的风险厌恶水平对系统的影响要比制造商的风险厌恶水平对系统的影响更显著。

参考文献:

- [1] 王志平,王众托.超网络理论及其应用[M].北京:科学出版社,2008:1-3.
- [2] 徐福缘,何静,林凤,等.多功能开放型企业供需网及其支持系统研究:国家自然科学基金项目(70072020)回溯[J].管理学报,2007,4(4):379-383.
- [3] 何建佳,徐福缘.面向和谐:管理的嬗变与多功能开放型企业供需网:基于复杂性、人类合作视角的分析[J].外国经济与管理,2009,31(3):16-22.
- [4] 徐福缘.大批量定制生产的理论与应用[M].上海:上海科学技术出版社,2008:1-4.
- [5] RAJU J,ZHANG Z J. Channel coordination in the presence of a dominant retailer[J]. *Marketing Science*,2005,24(2):254-262.
- [6] LAU A H L, LAU H S, WANG Jian-cai. Pricing and volume discounting for a dominant retailer with uncertain manufacturing cost information[J]. *European Journal of Operational Research*,2007,183(2):848-870.
- [7] HUA Zhong-sheng, LI Si-jie. Impacts of demand uncertainty on retailer's dominance and manufacturer-retailer supply chain coordination[J]. *Omega*,2008,36(5):697-714.
- [8] WANG C X, WEBSTER S. The loss-averse newsvendor problem[J]. *Omega*,2009,37(1):93-105.
- [9] CHOI S Y, RUSZCZYNSKI A. A risk-averse newsvendor with law invariant coherent measures of risk[J]. *Operations Research Letters*,2008,36(1):77-82.
- [10] WANG C X, WEBSTER S, SURESH N C. Would a risk averse newsvendor order less at a higher selling price? [J]. *European Journal of Operational Research*,2009,196(2):544-553.
- [11] PNG I P L, WANG Hao. Buyer uncertainty and two-part pricing: theory and applications[J]. *Management Science*,2010,56(2):334-342.
- [12] BERTSIMAS D, DOAN X V, NATARAJAN K, et al. Models for minimax stochastic linear optimization problems with risk aversion[J]. *Mathematics of Operations Research*,2010,35(3):580-602.
- [13] CHEN F Y, YANO C A. Improving supply chain performance and managing risk under weather-related demand uncertainty [J]. *Management Science*,2010,56(8):1380-1397.
- [14] 柳键,邱国斌,黄健.考虑缺货损失情形下损失厌恶零售商的订货决策[J].控制与决策,2012,27(8):1195-1200.
- [15] CACHON G P. Supply chain coordination with contracts[M]. [S. l.]: Elsevier,2003.
- [16] DENG Xiao-xue, XIE Jin-xing, XIONG Hua-chun. Manufacturer-retailer contracting with asymmetric information on retailer's degree of loss aversion[J]. *International Journal of Production Economics*,2013,142(2):372-380.

(上接第347页)

参考文献:

- [1] DUCHI J, SINGER Y. Online and batch learning using forward looking subgradients[C]//Proc of NIPS Workshop on Optimization for Machine Learning. 2008.
- [2] GAO P X, CURTIS A P, WONG B, et al. It's not easy being green [C]//Proc of ACM SIGCOMM Conference on Applications, Technologies, Architectures, and Protocols for Computer Communication. New York: ACM Press,2012:211-222.
- [3] LANGFORD J, LI Li-hong, ZHANG Tong. Sparse online learning via truncated gradient[J]. *Journal of Machine Learning Research*,2009,10(12):777-801.
- [4] DIETRICH K, LATORRE J M, OLMOS L, et al. Demand response in an isolated system with high wind integration[J]. *IEEE Trans on Power Systems*,2012,27(1):20-29.
- [5] MA J, SAUL L K, SAVAGE S, et al. Identifying suspicious URLs: an application of large-scale online learning[C]//Proc of the 26th Annual International Conference on Machine Learning. New York: ACM Press,2009:681-688.
- [6] HALKO N, MARTINSSON P, TROPP J. Finding structure with randomness: probabilistic algorithms for constructing approximate matrix decompositions[J]. *SIAM Review*,2011,53(2):217-288.
- [7] EMADANI O, BUI H, YE H E. Efficient online learning and prediction of users' desktop actions [C]//Proc of the 21st International Joint Conference on Artificial Intelligence. San Francisco: Morgan Kaufmann Publishers,2009:1457-1462.
- [8] HAN Yan-bo, WANG Gui-ling, JI Guang, et al. Situational data integration with data services and nested table[J]. *Service Oriented Computing and Application*,2012,7(2):129-150.
- [9] WANG Gang, GUI Xiao-lin. DRTEMBB: dynamic recommendation trust evaluation model based on bidding[J]. *Journal of Multimedia*,2012,7(4):279-288.
- [10] MANJON J V, COUPE P, BUADES A, et al. New methods for MRI denoising based on sparseness and self-similarity[J]. *Medical Image Analysis*,2012,16(1):18-27.
- [11] DREDZE M, CRAMMER K, PEREIRA F. Condence-weighted linear classification[C]//Proc of the 25th International Conference on Machine Learning. New York: ACM Press,2008:264-271.
- [12] RAUHUT H. Compressive sensing, structured random matrices and recovery of functions in high dimensions [C]//Proc of International Conference on Multivariate Approximation. 2011:1990-1993.
- [13] KIVINEN J, WARMUTH M K. Exponentiated gradient versus gradient descent for linear predictors [J]. *Information and Computation*,1997,132(2):1-63.