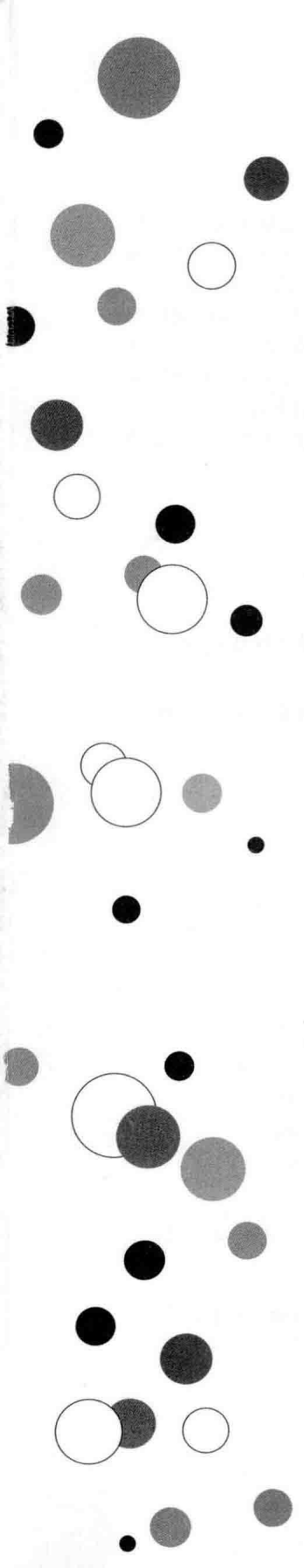


王建国 张文兴 等 著

支持向量机 建模及其智能优化

清华大学出版社



王建国 张文兴 等 著

支持向量机 建模及其智能优化

清华大学出版社
北京

内 容 简 介

本书以实际工业过程为背景,研究了冶金生产过程、化工生产过程,并获取了大量的实际生产数据,利用数据挖掘方法,挖掘出数据中隐含的生产规律,提出了一些改进的建模方法和优化方法,这些方法不仅适用于冶金和化工,还可广泛用于装备制造、材料、航天航空等领域。

本书的建模方法主要以支持向量机为基础,针对数据分布不平衡和海量数据的建模问题、模型的在线学习和优化问题进行了相关的算法研究,如粒度支持向量机、主动学习的增量支持向量机、误差校正的混合核函数在线支持向量机、粒子群智能优化方法和蚁群智能优化方法等,这些方法均配有仿真实验和实际生产数据实验,用于验证方法的有效性。

本书作为专业学术类参考书,可供高年级本科生、研究生、工程师、高校教师等人员阅读参考。

版权所有,侵权必究。侵权举报电话:010-62782989 13701121933

图书在版编目(CIP)数据

支持向量机建模及其智能优化/王建国,张文兴等著.--北京:清华大学出版社,2015

ISBN 978-7-302-40088-2

I. ①支… II. 王… ②张… III. ①向量计算机—系统建模—研究 IV. ①TP338

中国版本图书馆CIP数据核字(2015)第089642号

责任编辑:庄红权 洪 英

封面设计:常雪影

责任校对:王淑云

责任印制:沈 露

出版发行:清华大学出版社

网 址: <http://www.tup.com.cn>, <http://www.wqbook.com>

地 址:北京清华大学学研大厦A座 邮 编:100084

社总机:010-62770175 邮 购:010-62786544

投稿与读者服务:010-62776969, c-service@tup.tsinghua.edu.cn

质 量 反 馈:010-62772015, zhiliang@tup.tsinghua.edu.cn

印 刷 者:三河市君旺印务有限公司

装 订 者:三河市新茂装订有限公司

经 销:全国新华书店

开 本:170mm×240mm 印 张:11.5 字 数:213千字

版 次:2015年9月第1版 印 次:2015年9月第1次印刷

印 数:1~1000

定 价:39.00元

FOREWORD

20 世纪 90 年代, Vapnik 等人在 VC 维理论和结构风险最小化原则的基础上, 提出了针对小样本、非线性、高维问题的机器学习方法——支持向量机 (support vector machine, SVM)。历经十多年的发展, 支持向量机已成为国内外专家、学者的研究热点, 并在模式识别、光谱分析、数据挖掘、人工智能、医学、经济、社会等各个领域得到广泛应用。

多年来, 作者一直从事基于数据驱动的质量建模和群智能优化方面的研究, 特别是在支持向量机建模及粒子群和蚁群智能优化方法方面取得了一些有价值的研究成果。本书主要以支持向量机理论为基础, 针对数据分布不平衡和海量数据的 SVM 建模问题、模型的在线学习和优化问题以及相关的算法研究, 同时利用公开标准数据集和实际生产过程产生的实际数据进行了方法验证。

本书可供机械、信息、冶金、化工等领域的高年级本科生、研究生、工程技术人员和科技工作者阅读参考。本书具有如下特点。

(1) 着重从应用的角度出发, 强调理论与工程实际的紧密结合, 突出实用性。

(2) 内容由浅入深且易于理解。

(3) 每章均有典型的应用实例, 并给出了详实的步骤及其结果分析。

本书各章的主要内容包括: 第 1 章介绍统计学习理论、支持向量机的分类模型和回归模型、支持向量机的研究现状以及支持向量机的参数优化方法; 第 2 章分析支持向量机二次规划问题中的优化目标和约束条件问题, 介绍一种新的支持向量机快速建模方法——连续过松弛和严格凸二次规划的支持向量机, 并将该方法应用于冷轧带钢热镀锌生产过程建模; 第 3 章讨论数据分布不平衡情况下的建模方法, 介绍粒度计算和模糊 C 均值聚类方法, 在此基础上详细分析了模糊 C 均值和共享最近邻 (shared nearest neighbor, SNN) 相似度的粒度支持向

量机方法,最后将该方法用于甲醇合成过程的质量建模;第4章讨论大数据的增量建模方法,分析凸壳理论的几何意义及其求解方法,结合KKT条件,介绍基于凸壳和KKT条件的增量支持向量机建模方法;第5章讨论支持向量机在线建模方法,分析混合核函数及误差校正方法对在线建模的影响;第6章讨论粒子群智能优化算法,从粒子群的更新策略、惯性权重、多种群3个方面分析了粒子群方法的特点,并用粒子群算法优化支持向量机的结构参数;第7章讨论蚁群智能优化方法,分析蚁群优化方法的原理和特点,介绍分段蚁群优化方法和权重蚁群优化方法。

本书涉及的研究内容、成果得到了国家自然科学基金“基于支持向量机和群智能的煤制甲醇合成过程建模及优化方法研究”(21366017)和内蒙古自然科学基金重大项目“基于数据驱动的带钢热镀锌生产过程工艺参数优化的研究”(2011ZD08)资助。本书由王建国和张文兴完成主要撰写工作,此外,秦波参与了第2章和第3章、杨斌参与了第4章和第7章、刘文婧参与了第6章的撰写工作,云海滨、丛宽、范凯、陈良武、武丽明、赵元元、张志杰、陈肖洁等硕士研究生参与了部分研究工作,在此向他们表示感谢。

当今世界正处于知识大爆炸的时代,知识的更新日新月异,由于作者理论水平有限,以及所做研究工作的局限性,书中难免存在不妥之处,恳请广大读者批评指正。

作 者

2015年6月

CONTENTS

第 1 章 绪论	1
1.1 引言	1
1.2 模型预测方法	2
1.3 统计学习理论基础	3
1.3.1 VC 维	3
1.3.2 经验风险最小化原则	3
1.3.3 结构风险最小化原则	4
1.4 支持向量机的提出	7
1.5 支持向量机理论	8
1.5.1 分类支持向量机	8
1.5.2 回归支持向量机	12
1.6 支持向量机算法研究	16
1.6.1 块算法	16
1.6.2 分解算法	17
1.6.3 并行学习算法	17
1.6.4 原始空间中的学习算法	18
1.6.5 集成学习算法	18
1.6.6 复杂条件下的学习算法	18
1.7 支持向量机的参数优化	19
1.7.1 参数对支持向量机的影响	19
1.7.2 遗传算法优化支持向量机	20
1.7.3 蚁群算法优化支持向量机	21

1.7.4 粒子群算法优化支持向量机	22
参考文献	22
第2章 连续过松弛和严格凸二次规划的支持向量机	28
2.1 支持向量机的大规模训练样本方法	28
2.1.1 块算法	29
2.1.2 固定工作样本集法	29
2.1.3 增量学习算法	30
2.1.4 最小二乘算法	32
2.1.5 连续过松弛算法	33
2.2 连续过松弛和严格凸二次规划的支持向量机建模方法	34
2.2.1 严格的凸二次规划	34
2.2.2 快速支持向量机算法 FSVM	35
2.2.3 仿真实验及结果	37
2.2.4 小结	43
2.3 基于 FSVM 方法在带钢连续热镀锌质量建模中的应用	43
2.3.1 带钢连续热镀锌生产概述	43
2.3.2 冷轧热镀锌带钢的产品质量及其影响因素	43
2.3.3 模型参数的确定和样本的收集	45
2.3.4 数据预处理	46
2.3.5 预测模型的评判和参数选择	46
2.3.6 建立锌层重量预测模型	47
2.3.7 小结	49
参考文献	49
第3章 基于模糊 C 均值和 SNN 相似度的粒度支持向量机	52
3.1 粒度计算	52
3.1.1 词计算模型	53
3.1.2 粗糙集模型	54
3.1.3 商空间模型	55
3.2 粒度支持向量机	56
3.2.1 粒度支持向量机的研究现状	56
3.2.2 粒度支持向量机理论	57
3.3 核模糊 C 均值聚类	59

3.3.1	K 均值聚类	59
3.3.2	模糊 C 均值聚类	60
3.3.3	核模糊 C 均值聚类	62
3.4	基于核模糊 C 均值聚类的粒度支持向量机	63
3.4.1	算法原理及流程	63
3.4.2	仿真实验	64
3.4.3	小结	67
3.5	基于 SNN 相似度的粒度支持向量机	68
3.5.1	共享最近邻相似度	68
3.5.2	k 最近邻连通度	69
3.5.3	GSVM-SNN 算法步骤	71
3.5.4	仿真实验	72
3.5.5	小结	76
3.6	粒度支持向量回归机在甲醇合成中应用	77
3.6.1	粒度支持向量回归机	77
3.6.2	甲醇合成过程	78
3.6.3	影响甲醇合成的关键参数	80
3.6.4	甲醇合成建模	82
	参考文献	84
第 4 章	基于凸壳和 KKT 条件的增量支持向量机	87
4.1	支持向量机增量学习	87
4.2	主动学习	91
4.3	凸壳理论	92
4.4	基于 KKT 条件约束的增量支持向量机	95
4.4.1	算法原理	95
4.4.2	算法步骤	97
4.4.3	实验分析	97
4.5	基于凸壳和 KKT 条件约束的增量支持向量机建模方法	102
4.5.1	主动学习算法性能	103
4.5.2	基于主动学习的支持向量机增量学习算法性能	104
4.6	小结	108
	参考文献	109

第 5 章 误差校正的混合核函数在线支持向量机	112
5.1 在线支持向量机	112
5.2 精确在线支持向量回归机算法原理	114
5.3 核函数分析	119
5.3.1 单一核函数	119
5.3.2 混合核函数	120
5.4 混合核函数的精确在线支持向量机	121
5.4.1 数据预处理	121
5.4.2 算法步骤	122
5.4.3 仿真实验	123
5.4.4 小结	124
5.5 基于误差校正的混合核函数精确在线支持向量机	125
5.5.1 误差校正	125
5.5.2 仿真实验	126
5.5.3 小结	128
5.6 在线支持向量机在甲醇合成中的应用	128
5.6.1 模型的参数选择	128
5.6.2 甲醇合成在线支持向量机建模	129
参考文献	132
第 6 章 基于邻域自适应选取和双种群的粒子群优化支持向量机	133
6.1 粒子群优化方法	133
6.1.1 粒子群优化算法概述	133
6.1.2 标准粒子群算法	134
6.1.3 基本流程	135
6.2 非线性惯性权重和邻域自适应选取的粒子群优化支持向量机	136
6.2.1 对粒子速度与位置更新策略的改进	136
6.2.2 对惯性权重搜索方法的改进	137
6.2.3 IPSO 优化 SVM	137
6.2.4 实验分析	139
6.3 双种群的粒子群优化支持向量机	143
6.3.1 基于双种群的粒子群优化算法	143
6.3.2 双种群粒子群优化算法寻优过程模拟	145

6.3.3 基于 DP-PSO 优化 SVM 的步骤及流程	150
6.3.4 预测模型的评价指标	151
6.3.5 实验分析	152
6.3.6 小结	154
参考文献	155
第 7 章 权重分配的分段蚁群算法优化支持向量机	157
7.1 蚁群算法	157
7.1.1 蚁群算法的原理	157
7.1.2 蚁群算法的数学模型	158
7.1.3 蚁群算法的特点	160
7.1.4 蚁群算法的若干改进	161
7.2 分段蚁群算法优化 SVM	163
7.3 权重分配蚁群算法优化 SVM	165
7.4 基于权重分配的分段蚁群优化 SVM 的甲醇合成转化率预测 ...	170
7.5 小结	173
参考文献	173

第1章

绪 论

1.1 引言

人类智慧中一个重要的方面表现为从过去的数据和以往的知识中学习的能力。通过归纳学习,得到对客观世界的认识和规律。这些规律可以帮助人类认识现有世界,同时帮助人类对未知现象做出正确的预测和判断。人们把从过去的数据和以往的知识中学习并获取规律的能力称为学习能力。用获得的规律不仅可以解释已知实例,而且能够对未知现象做出正确的预测和判断,这种能力称为推广能力。

在人工智能的研究中,人们用计算机来模拟这种学习能力的问题常常被称为基于数据的机器学习问题。基于数据的机器学习的任务就是设计某种方法和模型,通过对已知数据的学习,找到数据内在的相互依赖关系,从而对已知数据进行预测和判断,其最终目的是使机器具有良好的推广能力。

统计学在机器学习中起着重要的基础作用,但是传统的统计学研究的主要是渐进理论,即当样本趋于无穷多时的统计性质。所以基于传统统计理论的各种学习方法,都是以样本无穷多为前提来推导算法。然而现实中,我们面对的样本数量往往是十分有限的,通常的方法是仍以样本无穷多为假设进行算法推导和建模。当样本数较少时,用这种方法得到的结果有时是难以令人满意的。例如,在普通的神经网络学习中,当样本数有限时,本来很好的学习机器却表现出很差的推广能力,这种现象被称为过学习现象。

为了解决这类问题,人们进行了坚持不懈的研究工作,直到 20 世纪 70 年代,Vapnik 等人开始建立一种研究有限样本情形下统计规律及学习方法性质的理论即统计学习理论(statistical learning theory, SLT)^[1,2],它为有限样本的机

器学习问题建立了一个良好的理论框架,较好地解决了小样本、非线性、高维数和局部极小点等实际问题,其核心思想就是学习机器要与有限的训练样本相适应。直到1995年,Vapnik和他的合作者们明确提出一种新的通用学习方法——支持向量机(support vector machine, SVM)^[1,2]后,该理论才受到广泛重视并应用到不同领域。支持向量机方法目前仍处于研究阶段,它仍存在许多有待进一步研究和探讨的问题,如为保证学习机器具有最佳推广能力而需进行模型参数的优化选择问题;在处理大规模数据和数据分布不平衡时,为提高模型的精度和速度问题;为了提高模型动态实时性的在线建模问题等。

1.2 模型预测方法

模型预测方法是机器学习的一个重要研究领域,绝大部分机器学习算法的应用都涉及对数据的预测性学习,目的是从已知数据中估计相关性以达到准确预测未来数据变化。经典的模式识别方法、各种统计方法以及近些年来出现的各种神经网络算法都被应用于估计这种数据中固有的相关性,试图建立一个能预测未来数据的模型。虽然,这些应用在一些特殊领域取得了一定进展,并且提出了一些新的概念和方法,但是,这些方法有其固有的算法缺陷:由于不知道数据的分布密度以及有限的训练数据,常常导致病态的训练结果,预测效果往往不尽如人意。人们迫切需要关于预测学习算法共同的概念和理论框架以及更高效的预测学习方法。到目前为止,几乎所有的预测学习算法都可以归结为以下三类。

(1) 传统的统计预测方法

这类方法是在参数结构形式已知的前提下,通过训练数据,预测各参数的值,如最小二乘法。应用这些预测方法除了需要很强的先验知识外,还需预先知道模型的结构形式。但是,在处理大量的实际预测问题时,常常不知道背景知识,面对一大堆采样数据,也不知道模型的结构形式。由于传统统计学研究的前提是样本数目趋于无穷大时的渐进理论,而参数预测方法几乎都是建立在这一前提基础之上的。因此只有当采样数据趋于无穷时,参数方法的训练结果才趋于真实模型。由于实际样本数目是有限的,很难满足这一前提。

(2) 经验非线性预测方法

20世纪80年代发展起来的人工神经网络和柔性统计方法就属于这一类。虽然,这些新方法克服了参数预测方法的部分弱点,能够依照需要,假设数据内在相关性而构造非线性模型。然而,这些非线性方法缺乏统一的数学理论基础,通常是从生物学的理论和一些学术流派中得到灵感,对于诸如神经网络中的结

构选择和权重初值的设定,仍需要借助于经验,得到的模型通常是局部最优解,而非全局最优。

(3) 统计学习理论

早期的统计学习理论是从 20 世纪 60 年代发展起来的,随着 VC 维理论和结构风险最小化原则的提出,进一步丰富和发展了统计学习理论,使它不仅是一种理论分析工具,还是一种能构造具有多维预测功能的预测学习算法的工具,使抽象的学习理论能够转化为通用的实际预测算法。

1.3 统计学习理论基础

统计学习理论被认为是目前针对小样本统计估计和预测学习的最佳理论。它从理论上系统地研究了经验风险最小化原则成立的条件、有限样本下经验风险与期望风险的关系,以及如何利用这些理论找到新的学习原则和方法等问题。

1.3.1 VC 维

统计学习理论是关于小样本进行归纳学习的理论,其中一个重要的概念就是 VC 维(Vapnik-Chervonenkis dimension)。模式识别方法中 VC 维的直观定义是: 对一个指示函数集,如果存在 h 个样本能够被函数集里的函数按照所有可能的 2^h 种形式分开,则称函数集能够把 h 个样本打散。函数集的 VC 维就是它能打散的最大样本数目 h 。若对任意数目的样本都有函数能将它们打散,则函数集的 VC 维是无穷大。有界实函数集的 VC 维可以通过用一定的阈值将它转化成指示函数集来定义。

VC 维反映了函数集的学习能力,一般而言,VC 维越大则学习机器越复杂,学习容量越大。目前尚没有通用的关于任意函数集 VC 维计算的理论,只对一些特殊的函数集知道其 VC 维。一般地,对于 n 维空间 $\mathbb{R}^n$ 中,最多只能有 n 个点是线性独立的,因此 $\mathbb{R}^n$ 空间超平面的 VC 维是 $n+1$ 。

但是对于非线性学习机器而言,VC 维与独立参数的个数之间并没有明确的对应关系,非但如此,在非线性情况下学习机器的 VC 维通常是无法计算的。但是实际中在应用统计学习理论时,可以通过变通的办法巧妙地避开直接求 VC 维的问题。

1.3.2 经验风险最小化原则

学习的目的是根据给定的训练样本求系统输入输出之间的依赖关系。学习问题可以一般地表示为变量 y 与 x 之间存在的未知依赖关系,即遵循某一未知

的联合概率 $F(\mathbf{x}, y)$ 。机器学习问题就是根据 n 个独立同分布观测样本

$$(\mathbf{x}_1, y_1), (\mathbf{x}_2, y_2), \dots, (\mathbf{x}_n, y_n) \quad (1.1)$$

在一组函数 $\{f(\mathbf{x}, \mathbf{w})\}$ 中求一个最优的函数 $f(\mathbf{x}, \mathbf{w})$ 对依赖关系进行估计, 使期望风险

$$R(\mathbf{w}) = \int L(y, f(\mathbf{x}, \mathbf{w})) dF(\mathbf{x}, y) \quad (1.2)$$

最小。其中, $\{f(\mathbf{x}, \mathbf{w})\}$ 称为预测函数集; $\mathbf{w}$ 为函数的广义参数。 $\{f(\mathbf{x}, \mathbf{w})\}$ 可以表示任何函数集, $L(y, f(\mathbf{x}, \mathbf{w}))$ 为由于用 $f(\mathbf{x}, \mathbf{w})$ 对 y 进行预测而造成的损失, 不同类型的学习问题有不同形式的损失函数。

在传统的学习方法中, 采用了所谓的经验风险最小化 (empirical risk minimization, ERM) 原则, 即用样本定义经验风险:

$$R_{\text{emp}}(\mathbf{w}) = \frac{1}{n} \sum_{i=1}^n L(y_i, f(\mathbf{x}_i, \mathbf{w})) \quad (1.3)$$

机器学习就是要设计学习算法使 $R_{\text{emp}}(\mathbf{w})$ 最小化。

在早期的神经网络研究中, 人们总是把注意力集中在如何使 $R_{\text{emp}}(\mathbf{w})$ 更小, 但很快发现一味追求训练误差小并不是总能达到好的预测效果。在某些情况下, 当训练误差过小反而会导致推广能力下降。此即几乎所有神经网络研究者都曾遇到过的过学习问题。出现过学习问题的原因有: 一是学习样本不充分, 二是学习机器设计不合理。这两个问题是相互关联的, 在神经网络学习中, 对于有限样本, 如果神经网络的学习能力过强, 足以记住每一个训练样本, 此时经验风险很快就可以收敛到很小甚至为零, 但是无法保证它对新训练样本能够得到好的预测。因此, 在有限样本情况下, 经验风险最小化并不一定意味着期望风险最小, 而且学习机器的复杂性不仅要与所研究的系统有关, 还要与有限的学习样本相适应。即有限样本下学习机器的复杂性和推广能力之间存在着矛盾, 采用复杂的学习机器容易使学习误差更小, 但却往往丧失了推广能力。

1.3.3 结构风险最小化原则

传统机器学习方法中普遍采用的经验风险最小化原则在样本数据有限时是不合理的, 因为经验风险和置信范围需同时最小。实际上, 在传统方法中, 选择学习模型和算法的过程就是优化置信范围的过程, 若选择的模型比较适合现有的训练样本 (相当于 h/n 值适当), 则可以取得比较好的结果。譬如在神经网络中, 需要根据问题和样本的具体情况来选择不同的网络结构 (对应不同的 VC 维), 然后进行经验风险最小化。

为解决经验风险和置信范围这两项最小化风险泛函问题, 引出结构风险最

小化(structural risk minimization, SRM)原则。

对指示函数集中的所有函数(包括使经验风险最小的函数),经验风险 $R_{\text{emp}}(\mathbf{w})$ 和实际风险 $R(\mathbf{w})$ 之间以至少 $1-\eta$ 的概率满足如下的关系:

$$R(\mathbf{w}) \leq R_{\text{emp}}(\mathbf{w}) + \sqrt{\frac{h[\ln(2n/h) + 1] - \ln(\eta/4)}{n}} \quad (1.4)$$

其中, h 是函数集的 VC 维; n 是样本数; η 是满足 $0 \leq \eta \leq 1$ 的参数。

由此可见,统计学习的实际风险 $R(\mathbf{w})$ 由两部分组成:一是经验风险(训练误差),另一部分称为置信范围(VC confidence)。置信范围反映了真实风险和经验风险差值的上界,反映了结构复杂所带来的风险,它和学习机器的 VC 维 h 及训练样本数 n 有关,式(1.4)可以简单地表示为

$$R(\mathbf{w}) \leq R_{\text{emp}}(\mathbf{w}) + \Phi(h/n) \quad (1.5)$$

由式(1.4)右边第二项可知, $\Phi(h/n)$ 随 h 加大而增长,因此,在有限训练样本下,学习机器的复杂性越高,VC 维越高,则置信范围越大,就会导致真实风险与经验风险之间可能的差别越大。

这里的泛化界限是对于最坏情况的结论,在很多情况下是较松的,尤其当 VC 维较高时更是如此。当 $h/n > 0.37$ 时,这个界限肯定是松弛的。当 VC 维无穷大时,这个界限就不再成立。

为了构造适合于小样本学习的归纳学习原理,可以通过控制学习机器的泛化能力来达到此目的。对于完全有界非负函数 $0 \leq f(\mathbf{x}, \mathbf{w}) \leq B$, 以至少 $1-\eta$ 的概率满足以下不等式:

$$R(\mathbf{w}) \leq R_{\text{emp}}(\mathbf{w}) + \frac{B\epsilon}{2} \left(1 + \sqrt{1 + \frac{4R_{\text{emp}}(\mathbf{w}_1)}{B\epsilon}} \right) \quad (1.6)$$

运用以上公式可以控制基于固定数目的经验样本之上的风险函数最小化过程。控制最小化过程有几种方法:一是最小化经验风险,根据上面的公式可以看出,风险的上界随着经验风险值的减少而减少;二是最小化式(1.6)中右边的第二项,将样本数目与 VC 维作为控制变量,后者适合于样本数较小的情况。

为了寻找使得结构风险达到最小的函数 $f(\mathbf{x}, \mathbf{w})$,统计学习理论提出了一种新的策略,它将函数集 $S = \{f(\mathbf{x}, \mathbf{w}), \mathbf{w} \in \Lambda\}$ 看成具有一定的结构,并由一系列嵌套的函数子集 $S_k = \{f(\mathbf{x}, \mathbf{w}), \mathbf{w} \in \Lambda_k\}$ 组成,它们满足

$$S_1 \subset S_2 \subset \cdots \subset S_k \subset \cdots \subset S_n$$

函数嵌套子集构成的函数集结构如图 1.1 所示。

为了选择合适的 S_k 作为学习函数,可以将式(1.4)右边划分为两部分,即第一项经验风险和第二项置信范围。如果给定样本数目 n ,那么,随着 VC 维数目 h 的增加,经验风险逐渐变小,而置信范围逐渐递增。如图 1.2 所示,真实风险

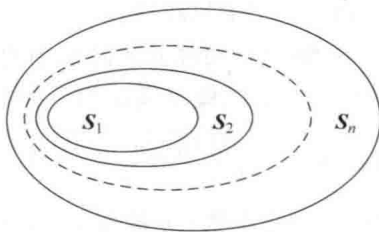


图 1.1 函数嵌套子集构成的函数集结构

的界是经验风险和置信范围之和,随着结构元素序号的增加,经验风险将减小,而置信范围将增加。最小的真实风险上界是在结构的某个适当的元素上取得的。综合考虑经验风险与置信区的变化,可以求得最小的风险边界,它所对应的函数集的中间子集 S^* 可以作为具有最佳泛化能力的函数集合。

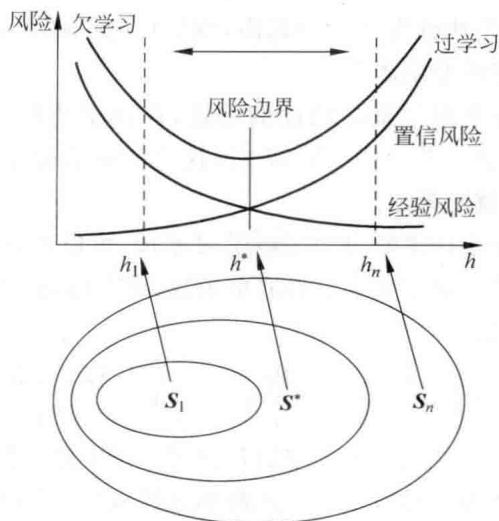


图 1.2 结构风险最小化原理图

根据图 1.2,可以得到以下两种运用结构风险最小化原理构造的学习机器的思想。

(1) 给定一个函数集合,按照上面的方法来组织一个嵌套的函数结构,在每个子集中求取最小经验风险,然后选择经验风险与置信风险之和最小的子集。当子集数目较大时,此方法较费时,甚至不可行。

(2) 构造函数集合的某种结构,使得在其中的各函数子集均可以取得最小的经验风险(如使得训练误差为 0)。然后,在这些子集中选择适当的子集使置信风险最小,则相应的函数子集中使得经验风险最小的函数就是所求解的最优

函数。

可见,利用结构风险最小化原理的思想,就可以完美解决传统机器学习中的过学习问题。支持向量机通过最大化分类边界以最小化 VC 维,也即在保证经验风险最小的基础上最小化置信风险,从而达到最小化结构风险的目的,因此支持向量机方法实际上就是结构风险最小化思想的具体实现。

1.4 支持向量机的提出

统计学习理论的 VC 维理论和结构风险最小化原则的提出为支持向量机算法打下了坚实的理论基石。

20 世纪六七十年代,非参数统计学以及不适定问题的解决方法的发现,特别是算法复杂度的概念及其与归纳推理的关系的提出,推动了统计学发展。算法复杂度的思想是统计学和信息论中最伟大的思想之一。到 20 世纪 70 年代,在算法复杂度思想的基础上,形成了对于学习问题的最小描述长度(minimum description length, MDL)归纳推理原则。如何在有限数量的经验数据基础上进行依赖关系估计的问题,这些思想改变了人们对算法复杂度的认识。1963 年, Vapnik、Lerner 以及 Chervonenkis 开始研究描述学习理论的非线性普遍算法,他们开创了实现统计机器学习算法的先河。从此,以 Vapnik 为代表的许多学者对统计机器学习的具体内容进行了不断地丰富和发展^[3]。

1968—1971 年, Vapnik 和 Chervonenkis 在《The Necessary and Sufficient Conditions for the Uniforms Convergence of Averages to Expected Values》一文中提出了一个重要的理论基础——VC 维理论^[4]。利用 VC 熵和 VC 维的概念,发现对指示函数空间的大数定律及其与模式识别问题的关系,并创造出对于经验风险最小化原则下的模式识别的一个一般的非渐近理论。

1974 年, Vapnik 形成结构风险最小化原则^[5], 1982 年, Vapnik 在《Estimation of Dependences Based on Empirical Data》^[6]一书中首次提出这一具有划时代意义的结构风险最小化原理,堪称为支持向量机算法的基石。

1974—1979 年, Vapnik 创造基于 ERM 和 SRM 原则上的一般的非渐近学习理论。

1981 年, Vapnik 把大数定律推广到实值函数空间^[4]。

1987 年, Vidyasaaga 出版了专著《A Theory of Learning and Generalization》,该书从数学角度详细讨论了统计学习理论。

1991 年, Vapnik 发现 ERM 原则与最大似然方法一致的充分必要条件^[7],完成了对经验风险最小化归纳推理的分析。