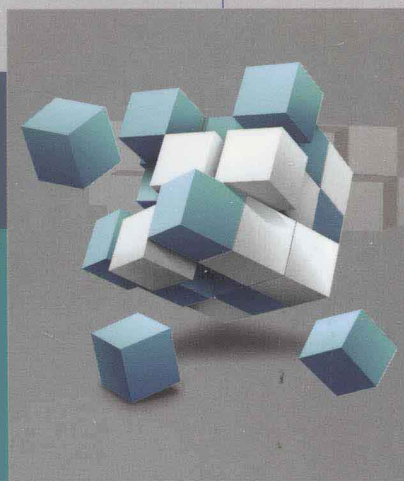


工业装备系统 亚健康 诊断方法

张 利 张立勇 王学芝 等著



电子工业出版社
PUBLISHING HOUSE OF ELECTRONICS INDUSTRY
<http://www.phei.com.cn>

本书由国家自然科学基金 (编号 : 61174115) 资助

工业装备系统亚健康诊断方法

张 利 张立勇 王学芝 等著

電子工業出版社

Publishing House of Electronics Industry

北京 • BEIJING

内 容 简 介

工业装备亚健康状态预测、诊断及维护技术能够使工业装备长期稳定运行,具有重要的现实意义。本书重点介绍了几种改进及新型工业装备健康状态诊断模型,并对如何将诊断模型应用到具体问题做了详细的阐述。全书共12章,第1章综述了各种机械健康状态诊断技术的发展现状及发展趋势;第2章主要介绍了一种有效的数据预处理方法;第3~10章给出几种改进及新型状态诊断模型;第11章和第12章介绍了不完整数据集的区间重构及在此基础上的聚类算法。

本书可作为高等院校机械工程、自动化等专业高年级本科生和研究生的学习参考书,也可供工程技术人员及研究人员参考使用。

未经许可,不得以任何方式复制或抄袭本书之部分或全部内容。
版权所有,侵权必究。



图书在版编目(CIP)数据

工业装备系统亚健康诊断方法/张利等著. —北京:电子工业出版社, 2013.11
ISBN 978-7-121-21714-2

I. ①工… II. ①张… III. ①工艺装备—维护 IV. ①TH16

中国版本图书馆 CIP 数据核字(2013)第 248020 号

策划编辑:秦绪军 徐蔷薇

责任编辑:王凌燕

印 刷:三河市双峰印刷装订有限公司

装 订:三河市双峰印刷装订有限公司

出版发行:电子工业出版社

北京市海淀区万寿路 173 信箱 邮编 100036

开 本:720×1000 1/16 印张:11.5 字数:232 千字

印 次:2013 年 11 月第 1 次印刷

定 价:39.00 元

凡所购买电子工业出版社图书有缺损问题,请向购买书店调换。若书店售缺,请与本社发行部联系,联系及邮购电话:(010) 88254888。

质量投诉请发邮件至 zlt@phei.com.cn, 盗版侵权举报请发邮件至 dbqq@phei.com.cn。

服务热线:(010) 88258888。

前言

随着科学技术的发展和进步,现代化生产中的工业系统变得越来越智能化,系统的复杂性成倍增长。近年来,我国的工业化水平越来越高,在航空航天、冶金化工、核能开发等领域取得了举世瞩目的成就。而目前国内外的工业装备正向着大规模和高精度的方向发展,工业装备中各零部件间衔接的紧密程度变得越来越重要。当其中的某个或某些零部件发生故障时,可能对整台设备乃至整条生产线造成重大影响,由此会带来无法估量的经济损失,更严重的还可能造成不必要的人员伤亡。因此,装备性能的下降和意外的失效影响生产中的三个关键因素,即安全、成本和质量。人们迫切需要提高动态系统的可靠性和可维护性。其中,如何在工业装备出现故障之前,准确监测工业装备的健康状况,判断出其亚健康状态,并进行调控、诊断和维护,以使其长期且稳定运行,具有重要的现实意义。

目前,国内有关工业装备亚健康状态预测、诊断及维护技术的书籍还较为少见,出版本书意在抛砖引玉,引起国内各界对工业装备监控维护的广泛关注。本书在对基础工业装备健康状态诊断理论介绍的基础上,结合作者近年来的研究成果,重点介绍了几种改进及新型工业装备健康状态诊断模型,并对如何将诊断模型应用到具体问题做了详细的阐述。全书共 12 章,第 1 章综述了各种机械健康状态诊断技术的发展现状及发展趋势;第 2 章主要介绍了一种有效的数据预处理方法;第 3~10 章给出几种改进及新型健康状态诊断模型,主要包括神经网络、支持向量机、小波变换和势能函数等几种方法;第 11 章和第 12 章介绍了不完整数据集的区间重构及在此基础上的聚类算法。

本书涉及的研究成果得到了国家自然科学基金项目(61174115, 51104044)和辽宁省教育厅项目(L2010153)的支持。本书在出版过程中,得到了很多学者的惠助。其中特别感谢吴利春、赵家强、杨永波、郑阿楠、田立、孙丽杰、邴兆

红、蔡晓辉、张旭男、夏天、王蓓蕾、刘萌萌、张继研、罗建华、蔡连博及卢秀颖等（以出现章节为序）参与了本书部分内容撰写或提供了素材。电子工业出版社徐编辑等为提高本书质量也付出了辛勤汗水，在此一并致谢。

需要指出的是，本书的部分内容是作者所在辽宁大学计算机应用研究所及大连理工大学计算机控制研究所师生多年从事模式识别、故障诊断工作的知识积累。在此向曾经或正在研究所工作学习、在故障诊断方面做出贡献的各位老师、同学和朋友表示感谢；同时，向给予作者支持和鼓励的清华大学、东北大学等兄弟院校的老师、朋友和同行表达谢意。

由于作者理论水平有限，书中难免存在疏漏与不妥之处，恳求各位专家、学者和广大读者批评指正。

作 者
2013 年 8 月

目 录

第 1 章 工业装备健康状态诊断方法论述	1
1.1 引言	1
1.2 粗糙集理论	2
1.2.1 粗糙集理论的相关概念	2
1.2.2 常用的属性约简算法	3
1.3 神经网络	5
1.3.1 BP 神经网络结构	5
1.3.2 BP 算法的步骤	6
1.3.3 BP 神经网络的性能分析	8
1.4 支持向量机	9
1.4.1 统计学习理论	9
1.4.2 支持向量机的理论	11
1.4.3 支持向量机的优点分析	13
1.5 小波分析	13
1.5.1 一维连续小波变换	14
1.5.2 离散小波变换	15
1.6 工业装备健康状态诊断	16
1.7 全书概况	17
参考文献	18
第 2 章 基于灰色粗糙集的二阶段数据预处理	22
2.1 引言	22
2.2 基于灰色粗糙集的二阶段数据预处理方法	22
2.2.1 关联度分析方法的基本理论	22
2.2.2 两阶段数据预处理算法流程	24
2.2.3 算法有效性验证	25
2.3 健康状态诊断中的特征参数提取	26
2.3.1 故障特征参数选取的原则	26

2.3.2 时域特征参数	27
2.3.3 频域特征参数	28
2.4 提取滚动轴承故障特征的仿真实验	29
2.4.1 仿真实验的故障数据	29
2.4.2 属性约简的实验过程	30
2.5 结束语	31
参考文献	31
第3章 基于遗传神经网络的健康状态诊断模型	33
3.1 引言	33
3.2 遗传算法与BP神经网络的结合	33
3.2.1 GA-BP结合的可行性分析	33
3.2.2 遗传算法与神经网络的结合方式	34
3.3 学习算子设计与改进	35
3.3.1 GA-BP编码方式	35
3.3.2 适应度函数的设计	36
3.3.3 选择算子的设计	37
3.3.4 交叉算子的设计	38
3.3.5 变异算子的设计	39
3.4 遗传神经网络健康状态诊断算法	40
3.4.1 算法的基本思想	40
3.4.2 算法的基本流程	41
3.5 健康状态诊断的仿真实验	42
3.5.1 仿真实验环境设置	42
3.5.2 对比实验与性能分析	43
3.6 结束语	46
参考文献	46
第4章 基于粒子群优化BP神经网络的亚健康识别	48
4.1 引言	48
4.2 粒子群算法概述	50
4.2.1 基本粒子群算法	50
4.2.2 带惯性权重的粒子群算法	50
4.2.3 带压缩因子的粒子群算法	51
4.3 粒子群算法的改进	52
4.3.1 精英学习策略的粒子群算法	52

4.3.2	算法改进的基本思想	53
4.3.3	惯性权重的改进	54
4.3.4	学习因子的改进	54
4.4	粒子群优化 BP 神经网络算法	55
4.4.1	IPSO-BP 模型	55
4.4.2	IPSO-BP 算法基本流程	57
4.5	亚健康及 D-S 证据理论的引入	58
4.5.1	亚健康	58
4.5.2	基于 D-S 证据理论的亚健康算法	59
4.6	健康状态诊断的仿真实验	62
4.6.1	实验过程	62
4.6.2	性能分析	67
4.7	结束语	67
	参考文献	68
第 5 章	基于机器学习的装备健康度评估	70
5.1	引言	70
5.2	模糊集理论	70
5.2.1	模糊集理论概述	70
5.2.2	相关概念	71
5.3	健康状态的等级划分	74
5.3.1	健康度的概念	74
5.3.2	故障状态的健康度评估	75
5.4	隶属度到健康度的映射关系模型	76
5.4.1	基于 BP 神经网络的健康度计算	77
5.4.2	基于支持向量机的健康度计算	78
5.5	健康状态诊断的仿真实验	80
5.5.1	故障特征参数灵敏度评估	80
5.5.2	健康度的计算	80
5.5.3	实验结果与分析	83
5.6	结束语	85
	参考文献	85
第 6 章	基于改进蚁群算法优化支持向量机参数的健康状态分类	88
6.1	引言	88
6.2	改进蚁群算法对支持向量机的优化过程	91

6.2.1 参数优化	91
6.2.2 基于网格划分策略的蚁群算法	92
6.3 仿真实验及结果分析	94
6.3.1 数据预处理	95
6.3.2 蚁群算法的参数设置	98
6.3.3 泛化能力	99
6.4 结束语	101
参考文献	102
第7章 基于概率神经网络的轴承健康状态诊断	105
7.1 引言	105
7.2 健康度定义	106
7.3 模型选择	107
7.3.1 概率神经网络概述	107
7.3.2 概率神经网络的拓扑结构	108
7.4 仿真实验及结果分析	109
7.5 结束语	115
参考文献	115
第8章 基于小波分析的健康状态检测	118
8.1 引言	118
8.2 脉冲小波	118
8.2.1 脉冲小波的定义	118
8.2.2 脉冲小波的正交性	119
8.3 脉冲小波分析方法	121
8.4 能量谱分析方法	122
8.5 健康状态诊断仿真实验	122
8.5.1 相关参数与实验结果	123
8.5.2 可行性分析	126
8.6 结束语	127
参考文献	127
第9章 势能函数健康状态识别分类算法的研究与应用	129
9.1 引言	129
9.2 势能函数	129
9.3 势能函数实现健康状态多分类	131

9.4 基于势能函数分类算法的健康状态诊断	132
9.5 结束语	134
参考文献	134
第 10 章 基于高斯混合模型 EM 算法的健康状态识别方法	136
10.1 引言	136
10.2 高斯混合模型的基本思想	136
10.2.1 高斯混合模型	136
10.2.2 GMM 的引入意义	137
10.2.3 EM 算法的改进思想	137
10.3 基于高斯混合模型 EM 算法的设计与实现	138
10.3.1 基于高斯混合模型 EM 算法	138
10.3.2 基于高斯混合模型 EM 算法的基本流程	139
10.4 仿真实验与结果	140
10.5 结束语	142
参考文献	142
第 11 章 不完整数据集的区间重构	144
11.1 引言	144
11.2 不完整数据集的处理	146
11.2.1 最近邻规则	147
11.2.2 不完整数据集转换为区间数据集	148
11.3 区间限定的必要性	149
11.4 不完整数据集的预分类	151
11.4.1 预分类方法分析	151
11.4.2 预分类过程	154
11.5 区间重构的流程	154
11.6 结束语	155
参考文献	155
第 12 章 基于区间重构的不完整数据集混杂聚类算法研究	158
12.1 引言	158
12.2 区间 FCM 聚类算法	158
12.2.1 区间模糊 C 均值基本算法	158
12.2.2 区间模糊 C 均值算法基本流程	159
12.2.3 基于最近邻区间的完整数据 FCM 算法	160

12.3 不完整数据集的粒子群模糊 C 均值混杂算法	162
12.3.1 群优化策略	162
12.3.2 粒子群的优点	162
12.3.3 混杂算法	163
12.3.4 混杂算法的变异	165
12.3.5 混杂算法基本流程	165
12.3.6 基于最近邻区间的不完整数据混杂聚类算法	166
12.4 基于区间重构的不完整数据集混杂算法	168
12.4.1 算法的基本思想	168
12.4.2 算法的基本流程	168
12.5 仿真案例及分析	170
12.6 结束语	171
参考文献	172

第 1 章 工业装备健康状态诊断方法论述

1.1 引言

改革开放以来，我国科学技术得到迅猛发展，工业化水平越来越高，现代化工业装备朝着大规模、高精度、复杂化过渡。装备各部件之间的联系、耦合更加紧密，当某一部件发生故障，整台设备、甚至是整条生产线都将受到影响，由此带来了无法计算的经济损失，更严重的还会产生不必要的人员伤亡。因此，复杂系统的可靠性和安全性迫切需要提高，以减少重大事故的发生。大量实践表明，健康状态诊断技术对保障设备的安全高效运行、尽早发现潜在故障、避免非计划停机、灾难性事故的发生起到至关重要的作用。因而，健康状态诊断评估技术成为一种用于解决现代化高精度系统装备可靠性、安全性的技术，是一种提高复杂设备安全使用的新方法。

对工业装备健康状态诊断方法的研究一直是一个热点，受到众多学者的关注，并取得一系列骄人的成果。清华大学信息学院的周东华教授从一个全新的视角，把现有的健康状态诊断方法分成定性分析和定量分析两大类，其中定量分析又分为基于解析模型的方法和基于数据驱动的方法，本书涉及的主要理论隶属于数据驱动方法的范畴，如图 1-1 所示^[1]。本章对本书涉及的理论的基本思想和研究进展

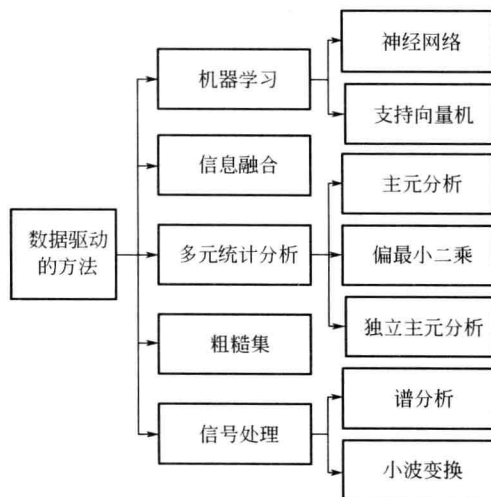


图 1-1 基于数据驱动的健康状态诊断方法分类

等做了较为详细的论述；此外，本章还概述了工业装备健康状态诊断的基本过程；最后，介绍了全书其他章节的主要内容。

1.2 粗糙集理论

粗糙集 (Rough Set) 是一种从数据中提炼知识并揭示暗含其中规律的数学工具，它能对各种不完备信息进行有效处理，并分析出隐藏的规律。属性约简是指在不影响整体决策分类能力的前提下，通过筛选重要或相关属性，用较少的重要属性总结出问题的正确决策。

粗糙集理论^[2]的最大特点在于，它不需要任何额外的相关数据信息，仅仅依赖已存于数据内部的知识，通过数据之间的近似来表达知识的不确定性，所以它对问题的不确定性的描述或处理是比较客观的。

1.2.1 粗糙集理论的相关概念

一般地，知识被认为是比较抽象的、具有普遍意义的理论，区别于传统知识的概念，粗糙集理论中的知识被看作是一种分类能力^[3~6]。

定义 1.1 由所研究对象组成的集合 U 称为论域，它是非空的、有限的。对于 U 上的一个非空子集 $\{X_1, X_2, \dots, X_n\}$ ，若 $X_i \cap X_j = \emptyset$ ， $i \neq j$ ，并且 $\bigcup_{i=1}^n X_i = U$ ，那么这个非空子集 $\{X_1, X_2, \dots, X_n\}$ 就是 U 上的一个划分。

定义 1.2 若集合 A 上的二元关系 R 满足自反关系、对称关系、传递关系，则称 R 是一个等价关系。对于 $\forall a \in A$ ， a 关于 R 的等价类 $[a]_R = \{x | xRa\}$ ，简称 $[a]_R$ 。用 U/R 表示，对于论域 U ，在等价关系 R 下的一个划分。

定义 1.3 知识系统 S 可以表示成有序的四元组 $S = \langle U, A, V, f \rangle$ 。其中，属性集 A 是由条件属性 C 和决策属性 D 的并集组成； $V = \bigcup_{r \in A} V_r$ 为属性值的集合， V_r 表示第 r 个属性的取值， $f: U \times A \rightarrow V$ 是一个映射函数，确定 U 中任意一个对象 x 的值。

定义 1.4 任意属性集 $\forall P \subseteq A$ ，知识系统 S 上的不可分辨关系 $\text{IND}(P)$ ，可定义为： $\text{IND}(P) = \{(x, y) \in U \times U : \forall b \in P, b(x) = b(y)\}$ 。

如果 $(x, y) \in \text{IND}(P)$ ，则称 x 和 y 对 P 是不可分辨的。显然，不可分辨关系 $\text{IND}(P)$ 为一等价关系，且 $\text{IND}(P) = \bigcap_{b \in P} \text{IND}(\{b\})$ 。 $U/\text{IND}(P)$ 是不可分辨关系 $\text{IND}(P)$ 对论域的一个划分，记作 U/P 。

定义 1.5 对于任意子集 $\forall X \subseteq U$ ， R 为不可分辨关系，给出 X 在知识系统 S

中关于 R 的下近似集合和上近似集合的定义:

$$\underline{R}(x) = \{Y \in U / \text{IND}(P) | Y \subset X\} \quad (1.1)$$

$$\bar{R}(x) = \{Y \in U / \text{IND}(P) | Y \cap X \neq \emptyset\} \quad (1.2)$$

称 $\text{POS}(X) = \underline{R}(x)$ 为正域, $\text{NEG}(X) = U - X$ 为负域, $\text{BN}(X) = \bar{R}(x) - \underline{R}(x)$ 为边界域。

由定义 1.5 可知, 在关系 R 下, 由那些肯定属于 X 中的元素构成了下近似集合 $\underline{R}(x)$, 又或者说是正域 $\text{POS}_R(X)$; 而由那些肯定不属于 X 的元素构成了负域 $\text{NEG}_R(X)$; 上近似集合 $\bar{R}(x)$ 是由那些肯定属于或可能属于 X 的元素组成的集合; 而边界域 $\text{BN}_R(X)$ 是那些不能确定的元素组成的集合。

在现实生活中, 使用众多特征属性来尽可能细致地描述对象, 这些属性的重要程度是不同的, 甚至说某些属性是冗余的。删除这些冗余的属性和不必要的属性值, 不但不会影响知识发现的结果, 还能达到简化计算的作用, 同时也能降低算法复杂度。

定义 1.6 设 R 是一等价关系集, 且 $r \in R$, 若有 $\text{IND}(R) = \text{IND}(R - \{r\})$ 成立, 那么 r 就是不必要的; 否则 r 是必要的。若对于 R 中的任何一个关系 $\forall r \in R$ 都是必要的, 那么称等价关系集 R 是独立的, 否则称 R 是依赖的。

定义 1.7 设 $Q \subseteq R$, 若 Q 是独立的, 且 $\text{IND}(Q) = \text{IND}(R)$, 则称 Q 是等价关系族 R 的一个约简, 记作 $\text{Red}(R)$ 。

定义 1.8 在等价关系 R 中, 由所有的必要关系构成的关系子集称为等价关系 R 的核, 记作 $\text{Core}(P)$ 。所有约简 $\text{Red}(R)$ 的交集等于 R 的核, 即 $\text{Core}(P) = \cap \text{Red}(R)$ 。

1.2.2 常用的属性约简算法

属性约简问题, 作为粗糙集理论的核心, 长久以来一直得到众多学者的关注, 先后出现众多行之有效的约简算法。在实际应用中, 没有必要获得属性的最佳约简, 只要得到一个有效的相对约简就已经足够可以有效解决具体问题了。

下面详细介绍两种常用的属性约简算法^[7-9]。

1. 基于属性重要度的属性约简

由核属性 $\text{Core}(P)$ 的定义可知, 各约简结果的交集一定属于核属性, 所以先找到 $\text{Core}(P)$, 有助于实现属性的快速约简。基于属性重要度的约简算法, 就是在求得属性集的核属性的基础上, 通过扩充得到属性子集的一种约简算法。

属性重要度计算公式为:

$$\text{SGF}(a, \text{Red}, X) = (|\text{POS}_{a \cup \text{Red}}(D)| - |\text{POS}_{\text{Red}}(D)|) / |U| \quad (1.3)$$

从非核属性集中找出部分属性, 然后将其与核属性合并所得并集计算对决策

属性的相对正域，当遍历完所有条件属性后，由相对正域相等的添加属性组成的集合就是一个约简结果。

具体实现步骤如下：

1) 首先求得属性集的核属性 Core

输入：信息系统 $S = (U, A, V, f)$ 。

输出：Core，表示核属性集。

(1) 加载数据 data，并令 $\text{Core} = \emptyset$ ；

(2) $\forall a \in A$ ，如果 $\text{IND}(A - \{a\}) \neq \text{IND}\{A\}$ ，则 $\text{Core} \leftarrow \text{Core}(Y\{a\})$ ；

(3) 当遍历完 A 中元素时，算法终止。

2) 属性约简

输入：信息系统 $S = (U, A, V, f)$ 。

输出：Red，是一个属性约简的结果。

(1) 输入数据 data，求出 S 的核属性集 Core；

(2) 令 $\text{Red} = \text{Core}$ ，若 $\text{IND}(\text{Red}) = \text{IND}(A)$ ，结束算法，输出 Red；

(3) 计算 $\forall a \in A - \text{Red}$ 的属性重要度并按降序排列，选取属性重要度最大的，当属性有多个时，则选取划分个数最少的， $\text{Red} \leftarrow \text{Red}(Y\{a\})$ ；

(4) 如果 $\text{IND}(\text{Red}) \neq \text{IND}(A)$ ，转到步骤 (3)，否则算法终止，此时的 Red 即为所有约简属性集。

2. 基于信息熵的属性约简

我们可以把等价关系认为，由论域 U 的任意子集构成的随机变量。

定义 1.9 等价关系 R 的信息熵 $H(R)$ 可以按下式来计算：

$$H(R) = - \sum_{i=1}^n P(X_i) \log P(X_i) \quad (1.4)$$

当 $H(R) = 0$ 时，说明只有一种可能性存在，没有其他不确定情况发生；相对于有 n 种可能发生的情况， $H(R)$ 取得越大，则系统的不确定性越大。

具体实现步骤如下：

输入：信息系统 $S = (U, A, V, f)$ 。

输出：Red，是一个属性约简后的结果。

(1) 按公式 (1.4) 计算 $H(A)$ 的值；

(2) 计算 S 的核属性集 Core，并在此基础上进行扩充，令 $\text{Red} = \text{Core}$ ；

(3) 当 $H(\text{Red}) \neq H(A)$ 时，顺序执行，否则跳转到步骤 (5)；

(4) 对于 $\forall a \in A - \text{Red}$ ，计算 a 的重要度：

$$H(\text{Red}(Y_a)) - H(\text{Red}) = \sum_{i=1}^{m_1} P(X_i) \log_2 P(X_i) + \sum_{i=1}^{m_2} P(Y_i) \log_2 P(Y_i) \quad (1.5)$$

其中, $\{P(X_i)\}_{i=1}^{m_1}$, $\{P(Y_i)\}_{i=1}^{m_2}$ 分别是 $\text{Red}(Y\{a\})$ 和 Red 的概率分布, 取差值最大的属性 a , 并令 $\text{Red} \leftarrow \text{Red}(Y\{a\})$, 跳转到步骤 (3);

(5) 算法结束, 输出 Red 。

1.3 神经网络

神经网络是一种机器学习算法, 而反向传播网络 (简称 BP 网络) 作为一种应用最广泛的神经网络, 具备一套完整的理论体系和成熟的学习机制。它模仿人脑神经元对外部激励信号的反应过程, 建立多层感知器模型, 利用信号正向传播和误差反向调节机制, 通过多次迭代学习, 搭建出处理非线性信息的网络模型。BP 网络及其算法是人工神经网络的精华所在, 它被广泛应用于函数逼近、模式识别、故障分类等。而在实际应用的人工神经网络, 80%~90% 都是直接采用 BP 网络或是其变种形式^[10]。

1.3.1 BP 神经网络结构

通常, BP 神经网络由输入层、隐含层和输出层构成, 相邻两层的节点彼此互相连接, 但位于同一层的节点间没有任何连接。隐含层节点个数目前尚无确定的标准, 可通过反复试凑的方法, 以相对较好的实验结果确定节点数。根据 Kolmogor 定理, 我们知道: 一个三层的 BP 神经网络, 如果隐含层采用 Signal 型的激活函数, 那么在闭区间上, 能以任意精度逼近相应的非线性连续函数。

为了简化说明, 本节选择具有单隐含层的 BP 神经网络进行说明, 即只有一层隐含层。如图 1-2 所示为一个具有单隐含层的 BP 神经网络模型^[11]。

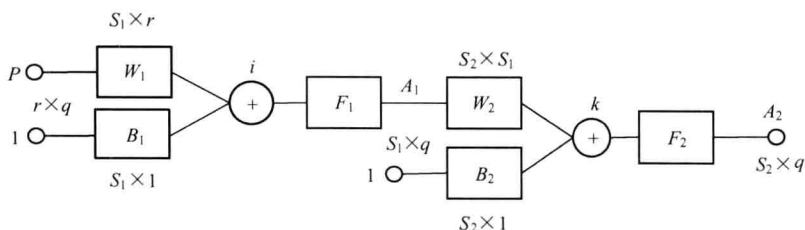


图 1-2 具有单隐含层的 BP 神经网络模型

在 BP 算法中, 各层神经元节点的激活函数必须是处处可导的, 因此处处可导的 S 形曲线被经常使用, 如对数函数和正切函数。在实际操作中, 纯线性函数常用作输出层的激活函数, 如图 1-3 所示为对数函数 $\text{logsig}()$ 和纯线性函数 $\text{purelin}()$ 。

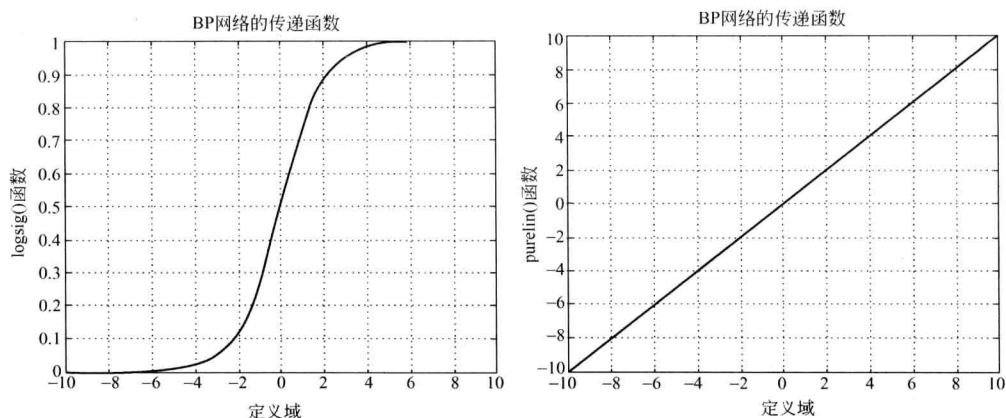


图 1-3 激活函数 logsig()和 purelin()

1.3.2 BP 算法的步骤

1. 标准 BP 算法的主要思想

标准 BP 算法^[12]把整个训练过程分成两个阶段，即信号的正向传播阶段和误差的反向调节阶段。在正向传播过程中，输入信号在权值和阈值的作用下经隐含层传向输出层，并在输出端产生累加计算结果。这一过程中，存在于网络中的连接权值和阈值保持不变，并且上一层神经元的输出只能作用于与其相邻的下一层神经元节点。当正向传播的结果与样本对的期望输出间的误差不能满足预先设定的误差精度时，转到误差的反向调节阶段。在误差的反向调整过程中，误差信号由输出端开始，按照学习算法向前传播，并将误差动态地分配给连接输出层与隐含层、隐含层与输入层的神经元间的权值。经过这种反复的调节，神经元间的权值和阈值得到不断的修正，并趋于稳定。在正向传播的计算结果与样本对的期望输出间的误差满足预先设定的误差精度或达到最大学习次数时，停止训练学习。

2. 标准 BP 算法的主要步骤

为了简化说明，本节以典型的三层 BP 网络为例，详细描述 BP 算法的计算流程，如图 1-4 所示。标准 BP 算法是一种基于梯度下降的学习算法，当网络输出与期望输出的误差达不到要求时，通过误差的反向调整过程，把误差分摊到存在于网络中的权重中，使正向传播的计算结果与样本对的期望输出间的均方误差趋于最小得以实现。