

结构健康监测： 理论建模和计算智能方法

JIEGOU JIANKANG JIANCE
LILUN JIANMO HE JISUAN
ZHINENG FANGFA

郑世杰 著



国防工业出版社
National Defense Industry Press

结构健康监测： 理论建模和计算智能方法

JIEGOU JIANKANG JIANCE
LILUN JIANMO HE JISUAN
ZHINENG FANGFA

郑世杰 著



国防工业出版社
National Defense Industry Press

内 容 简 介

本书是国内第一部系统论述智能算法在结构健康监测领域应用的著作。全书共分5章。第1章简要介绍结构健康监测技术所涉及的模糊推理、遗传算法等智能算法的基本概念、运算流程等。第2章到第5章是本书的核心。第2章侧重于研究基于遗传算法、遗传规划的非均匀应变分布重构方法;第3章对结构载荷识别进行论述,包括位于有限元网格节点上的混合编码载荷识别和位于单元内部的浮点数编码载荷识别以及动载荷的识别;第4章系统论述 RBF 神经网络、遗传算法、模糊聚类分析等的相互融合算法,并应用于复合材料结构的健康监测;第5章论述结构多损伤监测研究,包括基于传统遗传算法的多孔洞损伤和基于递阶遗传算法的任意程度多损伤监测。

希望本书能够为从事结构健康监测研究的科研人员和工程技术人员提供有益参考,本书也可作为高等院校相关专业的研究生教材。

图书在版编目(CIP)数据

结构健康监测:理论建模和计算智能方法/郑世杰
著. —北京:国防工业出版社,2014. 11
ISBN 978-7-118-09584-5

I. ①结… II. ①郑… III. ①工程结构—检测—研究
IV. ①TU3

中国版本图书馆 CIP 数据核字(2014)第 254091 号

※

国防工业出版社出版发行

(北京市海淀区紫竹院南路 23 号 邮政编码 100048)

北京嘉恒彩色印刷有限责任公司

新华书店经售

*

开本 710×1000 1/16 印张 13 $\frac{3}{4}$ 字数 239 千字
2014 年 11 月第 1 版第 1 次印刷 印数 1—2500 册 定价 36.00 元

(本书如有印装错误,我社负责调换)

国防书店:(010)88540777

发行传真:(010)88540755

发行邮购:(010)88540776

发行业务:(010)88540717

前言

近年来,科学家和工程师们从对自然界和生物体进化的学习与思考中得到启示,提出了一条力图从根本上解决工程结构在全生命周期内安全问题的崭新思路,即结构健康监测。结构健康监测技术是智能材料结构研究的一个重要方向,也是国内外研究的热点。计算智能则是人们受自然(生物界)规律的启迪,根据仿生学原理,模仿求解问题的算法,本书涉及的计算方法主要包括人工神经网络技术、遗传算法、遗传规划和模糊逻辑推理等。随着计算机技术的发展和普及,智能计算呈现相互融合的趋势,在最近 10 年得到了突飞猛进的发展,引起了诸多领域专家学者的关注,并成为结构健康监测领域一个跨学科的研究热点。运用计算智能解决结构健康监测过程中的一些问题,对相关学科研究人员既是一种严峻的技术挑战,也是一个难得的创新机遇。

南京航空航天大学机械结构及力学国家重点实验室是国内最早开展结构健康监测技术研究的单位之一,本书是在作者近 10 年来相关研究成果的基础上整理和发展而来,是国内第一部系统论述智能算法在结构健康监测领域应用的著作,研究内容遵循从简单到复杂的思路,研究对象从单损伤到多损伤,从离散的应变分布到连续的应变分布,从节点载荷识别到任意位置载荷识别,研究方法从单一智能算法的应用到多个智能算法的交叉融合。全书共分 5 章。第 1 章简要介绍了本书中结构健康监测技术所涉及到的模糊推理、遗传算法等智能算法的基本概念、运算流程等。第 2 章到第 5 章是本书的核心。第 2 章侧重于研究基于遗传算法的非均匀应变分布重构,包括重构算法、重构精度的影响因素分析、反射光谱的降噪回归等,还以寻求非均匀应变分布的表达式为出发点,将遗传规划方法应用于非均匀应变分布重构研究中,并进行了实验验证;第 3 章对结构载荷识别进行了论述,包括位于有限元网格节点上的混合编码载荷识别和位于单元内部的浮点数编码载荷识别以及动载荷的识别;第 4 章对用于结构健康监测的 RBF 神经网络及其与遗传算法、模糊聚类分析的交互融合进行了论述,构建了新的 RBF 神经网络训练算法,并用于复合材料结构脱层损伤监测,此外还研究了基于模糊逻辑推理的结构损伤监测方法;第 5 章论述了结构多损伤监测研

究,包括基于传统遗传算法的多孔洞损伤及基于递阶遗传算法的任意程度多损伤监测。

希望本书能够为从事结构健康监测研究的科研人员和工程技术人员提供有益参考,本书也可作为高等学校博士、硕士研究生的教材。

本书的出版得到了国防工业出版社的大力支持。感谢国家自然科学基金委,本书的很多工作得益于自然科学基金的支持。感谢我的同事王鑫伟教授、梁大开教授、袁慎芳教授多年的合作。感谢我的学生冯岩、张荣祥、李正强、夏彦君、李小平、宋振、范德礼、董会利、黄烨、王广帅、崔慧敏、张贵珍、张楠、王飞和马学仕等,本书的很多图表出自于他们之手,很多内容得益于他们的具体工作。

本书的研究内容涉及多学科交叉,限于作者水平,书中难免有不妥之处,恳请读者批评指正。

郑世杰

2014年4月

目录

第 1 章 几种主要的计算智能方法概论	1
1.1 模糊理论	1
1.1.1 模糊集合论基础知识	1
1.1.2 模糊关系	4
1.1.3 模糊推理	8
1.2 基本遗传算法	15
1.2.1 遗传算法的运算流程	15
1.2.2 遗传算法的编码	17
1.2.3 适应度函数	20
1.2.4 GA 的遗传操作	22
1.2.5 遗传算法的控制参数	29
1.2.6 约束条件的处理方法	30
1.3 遗传规划算法	31
1.3.1 函数的构成及表示	32
1.3.2 初始群体的生成	33
1.3.3 GP 的遗传操作	36
1.3.4 抑制规模爆炸	38
1.3.5 GP 的控制参数	39
1.3.6 遗传规划算法流程	40
1.4 计算智能与结构健康监测	41
第 2 章 FBG 轴向非均匀应变分布重构	43
2.1 反射谱的计算方法	45
2.1.1 非均匀 Bragg 光栅耦合模方程的龙格库塔解	45
2.1.2 非均匀 FBG 的 T 矩阵法分析	46
2.1.3 改进的传输矩阵法	47
2.2 非均匀应变分布重构问题描述与求解思路	50
2.3 混沌遗传算法	51

2.3.1	混沌优化	51
2.3.2	混沌遗传优化算法的构建	54
2.4	基于混沌遗传算法的非均匀应变分布重构研究	58
2.4.1	反射光谱采样	58
2.4.2	混沌比例选择	59
2.4.3	遗传算法参数设定	59
2.4.4	重构精度指标定义	61
2.4.5	应变分布重构的算例与结果分析	62
2.5	应变重构精度的影响因素分析	65
2.5.1	遗传代数对应变重构精度的影响	65
2.5.2	谱宽对应变重构精度的影响	65
2.5.3	谱采样间隔对应变重构精度的影响	66
2.6	基于样条插值的非均匀应变分布重构研究	69
2.7	基于遗传规划的非均匀应变分布重构算法	73
2.7.1	非均匀应变分布 GP 重构算法流程	74
2.7.2	遗传规划算法参数的设定	76
2.7.3	非均匀应变分布重构仿真分析	78
2.7.4	GP 与现有方法应变分布重构仿真效果的对比	82
2.8	基于支持向量回归的反射光谱去噪及应变分布重构精度研究	84
2.8.1	反射光谱测量噪声对重构精度的影响	84
2.8.2	用于反射光谱非线性回归的支持向量机	86
2.8.3	反射光谱降噪回归	88
2.8.4	基于支持向量回归反射光谱的应变分布重构算例分析	90
2.9	非均匀应变分布重构的实验验证	92

第3章 结构载荷识别 96

3.1	压电智能结构集聚电荷计算	97
3.2	RBF 神经网络在复合材料层合板、壳载荷识别中的应用	100
3.2.1	RBF 神经网络训练学习算法	101
3.2.2	RBF 神经网络样本数据的获取	103
3.2.3	RBF 神经网络载荷识别算例	103
3.3	基于混合编码遗传算法和有限元分析的压电结构载荷识别	106
3.3.1	用于结构载荷识别的适应度函数	107
3.3.2	混合编码遗传算法的编码方式	107

3.3.3	个体适应度的检测评估	107
3.3.4	混合编码遗传算法的遗传操作	108
3.3.5	载荷识别算法的流程设计	108
3.3.6	数值仿真算例	108
3.4	结构任意位置处的载荷识别	113
3.4.1	受载单元判别法	114
3.4.2	结构任意位置处载荷识别的浮点数编码方法	115
3.4.3	结构任意位置处载荷识别的算法流程	116
3.4.4	结构任意位置处载荷识别的仿真算例	116
3.5	结构动载荷识别的算法研究	119
3.5.1	有限元法进行动载荷识别的基本理论	119
3.5.2	动载荷识别的基本步骤与仿真算例	122
第4章	复合材料结构脱层损伤监测	129
4.1	基于混合递阶遗传 RBF 神经网络的复合材料结构脱层损伤 监测	130
4.1.1	递阶遗传算法	131
4.1.2	混合递阶遗传算法优化神经网络	133
4.1.3	混合递阶遗传算法设计	133
4.1.4	网络训练及损伤识别	136
4.2	基于遗传模糊 RBF 神经网络的复合材料结构脱层损伤监测	142
4.2.1	遗传模糊 RBF 神经网络结构	142
4.2.2	遗传模糊 RBF 神经网络的算法设计	144
4.2.3	基于遗传模糊 RBF 神经网络的复合材料脱层损伤网络 识别结果及分析	148
4.3	基于模糊逻辑的结构损伤监测	149
4.3.1	复合材料脱层损伤特征参数	150
4.3.2	基于模糊推理的损伤监测系统的建立	153
4.3.3	复合材料柔性梁脱层损伤识别	158
4.4	基于遗传模糊系统的结构损伤监测	164
4.4.1	遗传模糊系统简要介绍	165
4.4.2	遗传模糊系统的优化流程	166
4.4.3	仿真算例分析与结果	168

第5章 结构多损伤监测	172
5.1 基于遗传算法和拓扑优化的结构多孔洞损伤监测	173
5.1.1 孔洞损伤识别的拓扑优化公式	173
5.1.2 多孔洞损伤识别中的遗传算法设计	174
5.1.3 基于遗传算法和拓扑优化的多损伤识别研究实例仿真	177
5.2 梁结构多开口裂纹损伤监测	182
5.2.1 梁结构多开口裂纹损伤有限元模型	182
5.2.2 基于遗传算法和拓扑优化的梁结构开口裂纹多损伤监测方法	185
5.2.3 遗传算法改进策略与各种参数选择实例结果	186
5.3 基于递阶遗传算法的梁、板结构多损伤识别研究	191
5.3.1 针对悬臂梁多损伤监测的递阶编码	191
5.3.2 递阶遗传算法步骤与操作算子特点	192
5.3.3 基于递阶编码方法的悬臂梁多开口裂纹损伤监测仿真	193
5.3.4 递阶遗传算法对于二维板结构多损伤参数识别	198
5.3.5 多损伤板模型与损伤识别结果	199
附录 任意四边形受载单元受载点的等参坐标计算	203
参考文献	204

第1章

几种主要的计算智能方法概论

1.1 模糊理论

1.1.1 模糊集合论基础知识

1. 模糊集合的概念

在普通集合中,任何一个元素或个体与任何一个集合之间的关系只有“属于”和“不属于”两种情况,两者必居其一且仅居其一,绝对不允许模棱两可。然而,在现实世界中,有很多事物的分类边界是难以明确划分的,例如,将一群人划分为“高”和“不高”两类,就不好硬性规定一个划分的标准。如果硬性规定,1.80m 以上的人算“高个子”,否则不算,则有可能会出现两个本来身高“基本一样”的人却被认为一个“高”,一个“不高”,这就有悖于常理,因为这两个人在任何人看来都是“差不多高”。由此可见,普通集合在表达概念方面有它的局限性,普通集合只能表达“非此即彼”的概念,而不能表达“亦此亦彼”的现象。

为了对“属于”和“不属于”之间存在的无穷多层次、逐渐变化的灰度进行描述,1965 年美国加州大学控制专家 L. A. Zadeh 教授把普通集合扩展为模糊集合^[1, 2],提出用模糊集合来刻画模糊概念,这就是说,有些东西可以同时部分属于和部分不属于。

外延(或边界)不明确的集合称为模糊集合,由于模糊集合往往是某个论域的子集,所以常常称模糊集合为模糊子集,简称模糊集。本书中模糊集合都是用大写的英文字母下加波浪线表示。

2. 隶属度

对于上述模糊概念,不能仿照清晰概念用“属于”和“不属于”来描述,即不能简单地用普通集合中的特征函数来描述,必须通过反映某个元素隶属于模糊集合 $\underline{A}$ 的程度的隶属函数 $\mu_{\underline{A}}(x)$ 来描述。

定义 论域 U 中的模糊子集 $\underline{A}$, 是以隶属函数 $\mu_{\underline{A}}(x)$ 表征的集合, 即由映射

$$\begin{aligned}\mu_{\underline{A}}: U &\rightarrow [0, 1] \\ u &\rightarrow \mu_{\underline{A}}(u)\end{aligned}$$

确定论域 U 的一个模糊子集 $\underline{A}$ 。 $\mu_{\underline{A}}(x)$ 称为模糊子集 $\underline{A}$ 的隶属函数, 它表示论域 U 中的元素 u 属于其模糊子集 $\underline{A}$ 的程度。 $\mu_{\underline{A}}(x)$ 在 $[0, 1]$ 闭区间内可连续取值, $\mu_{\underline{A}}(x) = 1$, 表示 u 完全属于 $\underline{A}$; $\mu_{\underline{A}}(x) = 0$, 表示 u 完全不属于 $\underline{A}$; $0 < \mu_{\underline{A}}(x) < 1$, 表示 u 隶属于 $\underline{A}$ 的程度。

上述定义表明如下结果:

(1) 论域 U 中的元素是分明的, 即 U 本身是普通集合, 只是 U 的子集是模糊集合, 故称 $\underline{A}$ 为 U 的模糊子集。

(2) 隶属函数 $\mu_{\underline{A}}(x)$ 是用来说明 u 隶属于 $\underline{A}$ 的程度的, $\mu_{\underline{A}}(x)$ 的值越接近于 1, 表示 u 隶属 $\underline{A}$ 的程度越高。当 $\mu_{\underline{A}}(x)$ 的值域变为 $\{0, 1\}$ 时, 隶属函数 $\mu_{\underline{A}}(x)$ 蜕化为普通集合的特征函数, 模糊集合也就蜕化为普通集合, 所以说普通集是模糊集的特例。

(3) 正如普通集完全由特征函数刻画一样, 模糊集合完全由其隶属函数来刻画, 隶属函数是模糊数学的最基本概念, 借助于它才能对模糊集合进行量化。

3. 模糊集合的表示方法

1) Zadeh 表示方法

(1) 当 U 为离散有限论域 $U = \{u_1, u_2, \dots, u_n\}$ 时, 模糊集合 $\underline{A}$ 表示为

$$\underline{A} = \frac{\mu_{\underline{A}}(u_1)}{u_1} + \frac{\mu_{\underline{A}}(u_2)}{u_2} + \dots + \frac{\mu_{\underline{A}}(u_n)}{u_n} \quad (1.1)$$

式中, $\frac{\mu_{\underline{A}}(u_i)}{u_i}$ 不代表分式, 表示论域 U 中元素 u_i 及其隶属度 $\mu_{\underline{A}}(u_i)$ 之间的对应关系。符号“+”也不表示“加法”运算, 而是表示模糊集合在论域 U 上的整体, 是一种列举表示法。

(2) 当 U 为连续无限论域时, 模糊集合 $\underline{A}$ 表示为

$$\underline{A} = \int_u \frac{\mu_{\underline{A}}(u)}{u} \quad (1.2)$$

式中, 符号“ $\int$ ”不代表普通积分, 而是表示无限多个元素与其隶属度对应关系

的一个总括。

2) 向量表示法

当模糊集合 $\underline{A}$ 的论域由有限个元素构成时,模糊集合 $\underline{A}$ 表示成向量形式为

$$\underline{A} = [\mu_{\underline{A}}(u_1), \mu_{\underline{A}}(u_2), \dots, \mu_{\underline{A}}(u_n)] \quad (1.3)$$

注意:应用向量表示时,隶属度等于 0 的项不能舍弃,必须依次列入。

3) 序偶表示法

若将论域 U 中的元素 u_i 与其对应的隶属度 $\mu_{\underline{A}}(u_i)$ 组成序偶 $(u_i, \mu_{\underline{A}}(u_i))$, 模糊集合 $\underline{A}$ 可以表示为

$$\underline{A} = \{(u_1, \mu_{\underline{A}}(u_1)), (u_2, \mu_{\underline{A}}(u_2)), \dots, (u_n, \mu_{\underline{A}}(u_n))\} \quad (1.4)$$

4) 隶属函数法

用隶属函数的解析表达式表示出相应的模糊集合。

例如,Zadeh 曾用下述表达式来表达模糊集合“老年人”,其年龄论域 $U = [0, 100]$, u 为在论域上取值的年龄变量。

$$\mu_{\text{老年人}} = \begin{cases} 0, & 0 \leq u \leq 50 \\ 1 / \left[1 + \left(\frac{u - 50}{u} \right)^{-2} \right], & 50 < u \leq 100 \end{cases}$$

4. 模糊集合的基本运算

设 $\underline{A}$ 、 $\underline{B}$ 是论域 U 上的两个模糊集合,隶属函数分别为 $\mu_{\underline{A}}$ 和 $\mu_{\underline{B}}$,通常的运算有:

(1) “并”运算 $\underline{A} \cup \underline{B}$:

$$\underline{A} \cup \underline{B}(u) = \max[\mu_{\underline{A}}(u), \mu_{\underline{B}}(u)] = \mu_{\underline{A}}(u) \vee \mu_{\underline{B}}(u) \quad (1.5)$$

(2) “交”运算 $\underline{A} \cap \underline{B}$:

$$\underline{A} \cap \underline{B}(u) = \min[\mu_{\underline{A}}(u), \mu_{\underline{B}}(u)] = \mu_{\underline{A}}(u) \wedge \mu_{\underline{B}}(u) \quad (1.6)$$

(3) “补”运算 $\bar{\underline{A}}$:

$$\mu_{\bar{\underline{A}}}(u) = 1 - \mu_{\underline{A}}(u) \quad (1.7)$$

这里,符号“ $\vee$ ”“ $\wedge$ ”称为 Zadeh 算子,为模糊逻辑中的运算符号,在无限集合中,它们分别表示 $\sup$ 和 $\inf$;在有限元素之间,则表示 $\max$ 和 $\min$,即取最大值和最小值。

5. 隶属函数的确定

由于模糊集合理论研究的对象具有模糊性,再加上在客观实际工作中所研究的对象不同,因此,目前尚未有统一的隶属函数选择方法,但是仍有一些基本的原则可遵循,比如:隶属函数的选择要符合实际,不能违背常规知识。

隶属函数的确定大致有以下几种方法:

1) 模糊统计法

以调查统计所得结果,绘制出经验曲线作为隶属函数曲线,利用数学中曲线回归的方法,找出隶属函数的解析表达式。

2) 专家经验法

由专家的实际经验给出模糊信息的处理算式或相应权系数来确定隶属函数的方法。

3) 二元排序法

这是一种较实用的确定隶属函数的方法,它通过对多个事物之间两两对比来确定某种特征下的顺序,由此来决定这些事物对该特征的隶属函数的大致形状。

4) 典型函数法

根据问题的性质,应用一定的分析与推理,选用某些典型函数作为隶属函数。如三角形函数、梯形函数等。

1.1.2 模糊关系

1. 模糊关系的概念

客观世界中的各种事物间一般都存在某种联系,而描述客观事物间联系的数学模型就称作关系。普通关系 R 描述了事物之间“有”与“无”的肯定关系,但有些事物不能简单地用肯定或否定的词汇去明确表达它们之间的关系,例如家庭成员之间相貌的相似关系,就不是简单的相似或不相似,而是有着不同的相似程度,反映这种性质的关系就是模糊关系。集合论中的关系精确地描述了元素之间是否相关,而模糊集合论中的模糊关系则描述了元素之间相关的程度。模糊关系在概念上是普通关系的推广,普通关系则是模糊关系的特例。两个客体之间的关系称为二元关系,表示 3 个以上客体之间的关系称为多元关系。

2. 模糊关系矩阵

设 A 和 B 分别为论域 U 和 V 上的模糊集合,并且满足“若 A 则 B ”的关系,则称 $A \rightarrow B$ 是 $U \times V$ 上的模糊关系。对于模糊关系 $A \rightarrow B$ 的运算规则,有着很多的定义,这里介绍比较常用的 Zadeh 和 Mamdani 模糊关系定义方法。

1) 模糊关系的最大最小运算(Zadeh)

$$R = (A \times B) \cup (\bar{A} \times E) \quad (1.8)$$

式中, E 为全称矩阵,隶属函数形式的表达式为

$$\mu_{\tilde{R}}(x, y) = (\mu_{\tilde{A}}(x)) \wedge (\mu_{\tilde{B}}(y)) \vee (1 - \mu_{\tilde{A}}(x)) \quad (1.9)$$

2) 模糊关系的最小运算(Mamdani)

$$\tilde{R} = \tilde{A} \times \tilde{B} \quad (1.10)$$

隶属函数形式的表达式为

$$\mu_{\tilde{R}}(x, y) = (\mu_{\tilde{A}}(x)) \wedge (\mu_{\tilde{B}}(y)) \quad (1.11)$$

3) 模糊合成关系运算

设模糊矩阵 $\tilde{R} = (r_{ij})_{m \times n}$, $\tilde{S} = (s_{jk})_{n \times l}$ ($i = 1, 2, \dots, m, j = 1, 2, \dots, n, k = 1, 2, \dots, l$) 分别是 U 到 V 、 V 到 W 的模糊关系, 则定义 $\tilde{R}$ 、 $\tilde{S}$ 的合成 $\tilde{R} \circ \tilde{S}$ 是 U 到 W 的一个模糊关系, 称 $\tilde{T} = \tilde{R} \circ \tilde{S}$ 为模糊合成关系运算。因模糊关系合成采用的运算形式不同, 常见的有 4 种不同的方法: 最大-最小合成法, 最大-代数积合成法, 最大-有界积合成法和最大-强制积合成法, 其中 Zadeh 给出的最大-最小合成运算是最常用的模糊关系合成运算方法, 其隶属函数形式为

$$t_{ik} = \bigvee_{j=1}^n (r_{ij} \wedge s_{jk}) \quad (1.12)$$

模糊关系矩阵运算的基本规律与模糊集合运算的基本规律类同, 就不再作详细的介绍了, 若将“ $\vee$ ”换为“+”, 将“ $\wedge$ ”换为“ $\times$ ”, 则模糊矩阵的合成与普通矩阵的乘积完全相同。

3. 模糊向量的笛卡儿积

定义: 设两个模糊行向量 $\tilde{A}$ 和 $\tilde{B}$, 它们的笛卡儿积定义如下:

$$\tilde{A} \times \tilde{B} = \tilde{A}^T \circ \tilde{B} \quad (1.13)$$

4. 模糊聚类分析

聚类就是按照一定的要求和规律对事物进行区分和分类的过程, 在这一过程中仅靠事物间的相似性作为类属划分的准则。聚类分析则是指利用数学的方法研究和处理给定对象的分类^[3]。传统的聚类分析是一种硬划分, 它把每个待辨识的对象严格地划分到某类中, 具有非此即彼的性质, 因此这种类别划分的界限是分明的。而实际上大多数对象并没有严格的属性, 它们在性态和类属方面具有亦此亦彼的性质, 因此适合进行软划分。模糊集理论的提出为这种软划分提供了有力的分析工具, 人们开始用模糊的方法来处理聚类问题, 并称之为模糊聚类分析。

模糊 C 均值聚类(FCM)算法^[4]是目前最受欢迎和应用最广泛的聚类分析方法之一, 是用隶属度确定每个数据点属于某个聚类的程度的一种聚类算法, 1973 年, Bezdek 提出了该算法。FCM 是普通 C 均值聚类算法(即 K 均值聚类算法)的改进, 普通 C 均值算法对于数据的划分是硬性的, 而 FCM 则是一种柔性的模糊划分, 因此介绍 FCM 算法之前, 先引入 K 均值聚类算法。

1) K 均值聚类算法

K 均值聚类算法的核心思想如下: 算法把 n 个向量 x_j ($1, 2, \dots, n$) 分为 c 个

组 $G_i (i=1, 2, \dots, c)$, 并求每组的聚类中心, 使得非相似性(或距离)指标的价值函数(或目标函数)达到最小。当选择欧几里德距离为组 i 中向量 x_k 与相应聚类中心 c_i 间的非相似性指标时, 价值函数可定义为:

$$J = \sum_{i=1}^c J_i = \sum_{i=1}^c \left(\sum_{k, x_k \in G_i} \|x_k - c_i\|^2 \right) \quad (1.14)$$

这里 $J_i = \sum_{k, x_k \in G_i} \|x_k - c_i\|^2$ 是组 i 内的价值函数, 不难看出 J_i 的值依赖于 G_i 的几何特性和 c_i 的位置。

一般来说, 可用一个通用距离函数 $d(x_k, c_i)$ 代替组 i 中的向量 x_k , 则相应的总价值函数可表示为

$$J = \sum_{i=1}^c J_i = \sum_{i=1}^c \left(\sum_{k, x_k \in G_i} d(x_k - c_i) \right) \quad (1.15)$$

为简单起见, 这里用欧几里德距离作为向量的非相似性指标, 且总的价值函数表示为式(1.15)。

划分过的组一般用一个 $c \times n$ 的二维隶属矩阵 U 来定义, 如果第 j 个数据点 x_j 属于组 i , 则 U 中的元素 u_{ij} 为 1; 否则, 该元素取 0。一旦确定聚类中心 c_i , 可导出使式(1.15)最小的 u_{ij} 如下:

$$u_{ij} = \begin{cases} 1 & \text{对每个 } k \neq i, \text{ 如果 } \|x_j - c_i\|^2 \leq \|x_j - c_k\|^2 \\ 0 & \text{其他} \end{cases} \quad (1.16)$$

重申一点, 如果 c_i 是 x_j 的最近的聚类中心, 那么 x_j 属于组 i 。由于一个给定数据只能属于一个组, 所以隶属矩阵 U 具有如下性质:

$$\sum_{i=1}^c u_{ij} = 1, \quad \forall j = 1, \dots, n \quad (1.17)$$

且

$$\sum_{i=1}^c \sum_{j=1}^n u_{ij} = n \quad (1.18)$$

另一方面, 如果固定 u_{ij} , 则使(1.14)式最小的最佳聚类中心就是组 i 中所有向量的均值:

$$c_i = \frac{1}{|G_i|} \sum_{k, x_k \in G_i} x_k \quad (1.19)$$

这里 $|G_i|$ 是 G_i 的模或 $|G_i| = \sum_{j=1}^n u_{ij}$ 。

为便于批模式运行, 这里给出数据集 $x_i (1, 2, \dots, n)$ 的 K 均值算法, 该算法重复使用下列步骤, 确定聚类中心 c_i 和隶属矩阵 U :

步骤 1: 初始化聚类中心 $c_i, i=1, \dots, c$, 典型的做法是从所有数据点中任取 c 个点;

步骤 2: 用式 (1.16) 确定隶属矩阵 U ;

步骤 3: 根据式 (1.14) 计算价值函数, 如果它小于某个确定的阈值, 或它相对上次价值函数值的改变量小于某个阈值, 则算法停止;

步骤 4: 根据式 (1.19) 修正聚类中心, 返回步骤 2。

该算法本身是迭代的, 不能确保它收敛于最优解。K 均值算法的性能依赖于聚类中心的初始位置。所以, 为了使它可取, 要么用一些前端方法求好的初始聚类中心, 要么每次用不同的初始聚类中心, 将该算法运行多次。此外, 上述算法仅仅是一种具有代表性的方法, 还可以先初始化一个任意的隶属矩阵, 然后再执行迭代过程。

K 均值算法也可以在线方式运行, 这时, 通过时间平均, 导出相应的聚类中心和相应的组, 即对于给定的数据点 x , 用该算法求最近的聚类中心 c_i , 并用下面公式进行修正:

$$\Delta c_i = \eta(x - c_i) \quad (1.20)$$

2) 模糊 C 均值聚类

模糊 C 均值聚类(FCM)算法把 n 个向量 $x_i (i=1, 2, \dots, n)$ 分为 c 个模糊组, 并求每组的聚类中心, 使得非相似性指标的价值函数达到最小。FCM 与 HCM 的主要区别在于 FCM 用模糊划分, 使得每个给定数据点用值在 $[0, 1]$ 间的隶属度来确定其属于各个组的程度。与引入模糊划分相适应, 隶属矩阵 U 允许有取值在 $[0, 1]$ 间的元素。不过, 加上归一化规定, 一个数据集的隶属度的和总等于 1:

$$\sum_{i=1}^c u_{ij} = 1, \quad \forall j = 1, \dots, n \quad (1.21)$$

那么, FCM 的价值函数就是式 (1.14) 的一般化形式:

$$J(U, c_1, \dots, c_c) = \sum_{i=1}^c J_i = \sum_{i=1}^c \sum_{j=1}^n u_{ij}^m d_{ij}^2 \quad (1.22)$$

这里 u_{ij} 介于 $[0, 1]$ 间; c_i 为模糊组 i 的聚类中心, $d_{ij} = \|c_i - x_j\|$ 为第 i 个聚类中心与第 j 个数据点间的欧几里德距离, 且 $m \in [1, \infty)$ 是一个加权指数。

构造如下新的价值函数, 可求得使式 (1.22) 达到最小值的必要条件:

$$\begin{aligned} \bar{J}(U, c_1, \dots, c_c, \lambda_1, \dots, \lambda_n) &= J(U, c_1, \dots, c_c) + \sum_{j=1}^n \lambda_j \left(\sum_{i=1}^c u_{ij} - 1 \right) \\ &= \sum_{i=1}^c \sum_{j=1}^n u_{ij}^m d_{ij}^2 + \sum_{j=1}^n \lambda_j \left(\sum_{i=1}^c u_{ij} - 1 \right) \end{aligned} \quad (1.23)$$

这里 $\lambda_j (j=1, \dots, n)$ 是式 (1.21) 的 n 个约束式的拉格朗日乘子, 对所有输入参量求导, 使式 (1.22) 达到最小的必要条件为

$$c_i = \frac{\sum_{j=1}^n u_{ij}^m x_j}{\sum_{j=1}^n u_{ij}^m} \quad (1.24)$$

和

$$u_{ij} = \frac{1}{\sum_{k=1}^c \left(\frac{d_{ij}}{d_{kj}} \right)^{2/(m-1)}} \quad (1.25)$$

由上述两个必要条件, 模糊 C 均值聚类算法是一个简单的迭代过程, 在批处理方式运行时, FCM 用下列步骤确定聚类中心 c_i 和隶属矩阵 U :

步骤 1: 用值在 $[0, 1]$ 间的随机数初始化隶属矩阵 U , 使其满足式 (1.21) 中的约束条件;

步骤 2: 用式 (1.24) 计算 c 个聚类中心 $c_i, i=1, \dots, c$;

步骤 3: 根据式 (1.22) 计算价值函数, 如果它小于某个确定的阈值, 或它相对上次价值函数值的改变量小于某个阈值, 则算法停止;

步骤 4: 用式 (1.25) 计算新的 U 矩阵, 返回步骤 2。

上述算法也可以先初始化聚类中心, 然后再执行迭代过程。由于不能确保 FCM 收敛于一个最优解, 算法的性能依赖于初始聚类中心。因此, 要么用另外的快速算法确定初始聚类中心, 要么每次用不同的初始聚类中心启动该算法, 多次运行 FCM。

1.1.3 模糊推理

1. 模糊语言与语言变量

语言是一种以文字为符号的符号系统, 是人们表达思维的一种形式, 目前有两种大的语言类型——自然语言和人工语言。自然语言是人们在日常生活和工作中进行交流所使用的语言, 通过它描述了主、客观世界的各种事物、概念、行为、情感以及相互间的关系。自然语言突出特点是具有模糊性, 带有模糊性的语言, 称为模糊语言, 它的模糊性主要体现在语音、语义、语法等方面。模糊语言虽然不够精确和严格, 但它的这种模糊性使得自然语言更加生动和富有表现力, 也才显示出了人们判断和处理模糊现象的能力。

模糊语言变量的概念是 Zadeh 在 1975 年首先提出的, 所谓语言变量是以自然或人工语言中的字或句作为变量, 而不是以数值作为变量, 语言变量的值称为