

矿山智能选冶过程中多源 异构信息融合处理

5.1 智能矿山选冶综述

本章首先阐述了当前矿山智能选冶技术的现状及不足,随后对基于多源异构信息融合处理的矿山选冶技术进行概述。最后,以智能选冶技术的快速发展为背景,强调了多源异构信息融合处理在整个系统中的核心作用,分析了这些技术如何提升选冶过程的智能化与安全性。

5.1.1 矿山智能选冶的现状及不足

随着计算机技术和自动控制技术的进步,选冶过程的自动化和智能化水平不断提高,精细化管控成为智能选冶控制的主流。我国矿山智能选冶的工业背景是在全球工业革命的大背景下,通过技术融合、国家政策的推动以及建设水平的不断提升,逐步发展起来的。这一过程不仅反映了矿业技术的进步,也体现了国家对于矿业现代化和智能化转型的重视。

虽然我国矿山信息化建设取得了显著成绩,但仍存在许多问题。首先,智能矿山的概念在一些地区还不够清晰,对于智能矿山的理解还停留在数字化、自动化、无人化作业的层面,缺乏一个全面和深入的认识。其次,我国矿物赋存条件复杂多样,不同矿区的智能化建设基础不平衡。再次,智能选冶技术的发展还面临着技术和安全保障方面的问题。最后,智能选冶技术需要处理多传感器收集到的复杂数据,如何将这些复杂数据进行整合并充分利用来使得智能选冶过程更加高效准确也是目前智能选冶技术面临的一个难题。

5.1.2 信息融合在智能选冶过程中的重要作用

信息融合可以定义为利用计算机技术对按时序获得的若干传感器的观测信息在一定准则下进行自动分析、综合,以完成所需的决策和估计任务的信息处理过程。通过对定义进行分析可以得出,信息融合是以多传感器为基础,通过分级、分层的处理、分析等来对目标进行综合判断的过程,以此来提高对环境描述的准确性。信息融合的基本过程如图 5-1 所示。

综上所述,信息融合在智能选冶发展过程中扮演着至关重要的角色。信息融合技术的引入不仅能够将传感器充分利用到智能选冶过程中,提高矿山智能选冶的效率和效果,实现大数据背景下的工业精细化管控,还能够帮助企业应对环境挑战,减少选冶过程中产生的污染物,实现可持续发展。随着技术的不断进步,信息融合将在智能选冶领域发挥越来越重要的作用。

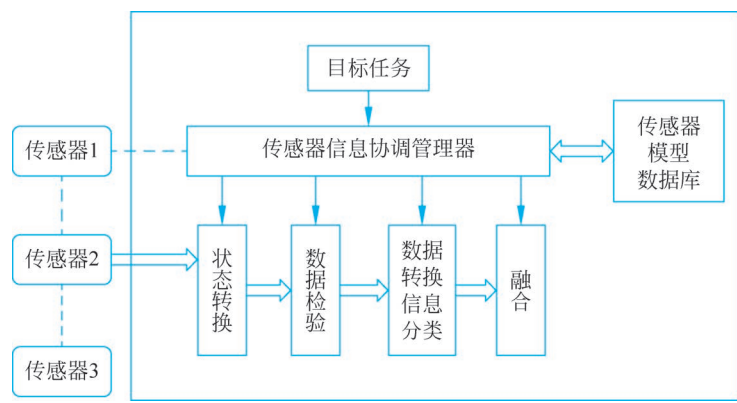


图 5-1 信息融合的基本过程

5.2 智能选冶多源异构信息预处理

在选冶场景下,由于信息的多源异构性、信息来源途径的多样化、矿山环境等,单来源的传感器信息难以完整表述和掌控矿山态势,但来源于不同传感器的数据充斥着不精确、不一致甚至可能有互相矛盾的信息。因此,多传感器信息融合技术应运而生并开始蓬勃发展,它使得多个传感器协同工作,把已有信息融合在一起,得到比单源传感器更准确可靠的信息。于是,多传感器数据融合这一研究方向能够迅速地发展起来并在智慧矿山中达到了较为普遍的应用。

5.2.1 信息融合数据预处理简介

信息融合数据预处理是对多源数据进行清洗、转换和整合的过程,以便为后续的融合、分析和决策提供准确、一致的数据集。数据预处理流程如图 5-2 所示。

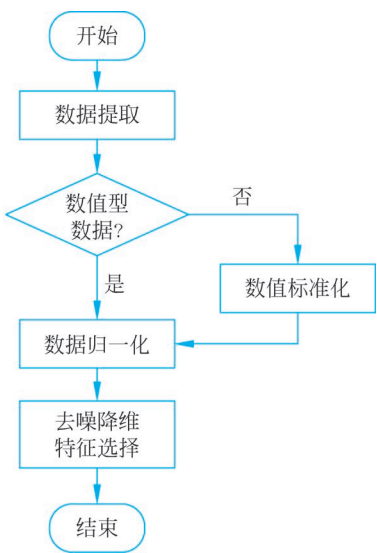


图 5-2 数据预处理流程

数据预处理在信息融合中起到十分重要的作用:这一过程可以消除原始数据中可能存在的噪声、错误和不一致性,识别并处理异常值,避免这些值对数据分析结果产生误导,从而提高数据质量,使数据更加准确和可靠;通过预处理,数据变得更加易于理解,增强数据可解释性,有助于分析师或决策者更好地解读数据,从而做出更加合理的判断和决策;在多源信息融合中,不同数据源可能有不同的格式或标准,数据预处理可以将这些数据统一到一致的格式或标准下,解决数据之间的对应关系和冲突问题,将来自不同来源的数据集成整合到一起,以便进行有效的融合分析;预处理后的数据可以减少分析过程中的错误和迭代次数,降低模型训练难度,提升整个数据分析流程的效率。部分数据预处理解决思路 and 具体方法如表 5-1 所示。

表 5-1 数据预处理解决思路和具体方法

问 题	思 路	情 况 分 析	具 体 方 法
缺失值处理	填充缺失数据	通过前后数据补全	使用前后数据的均值补全时间序列的缺失值
		缺失值过多	使用平滑等处理,如滑动窗口
		设置默认值	数据缺失时用默认值填充
		无法补全	剔除数据,但保留不删除
数据去重	去除重复记录	按主键去重	用 SQL 语言或脚本语言“去除重复记录”即可
		按规则去重	编写一系列的规则,对重复情况复杂的数据进行去重
噪声过滤	去除异常值	通过机器学习算法去噪	基于统计、基于邻近性、基于可视化、基于分类以及基于聚类
数据集成	具体问题具体分析	度量单位不一致	建立数据体系、统一维度单位等
		维度不一致	降维方法:主成分分析法、随机森林等 抽象方法:汇总、离散化等
		数据差异大	归一化:Min-Max、z-score 等

数据预处理是所有数据融合中必不可少的一步,预处理的结果作为数据融合的数据源,其质量将直接影响数据融合的结果,一个好的预处理结果不仅能够使融合的结果更加准确,还可以提高融合速度。

5.2.2 智能选冶数据预处理技术

智能选冶数据预处理技术是一种基于人工智能和机器学习算法的数据预处理技术,数据处理的工作时间占据了整个数据分析项目的 70% 以上。在对传感器收集到的信息进行预处理的过程中,为了确保数据的完整性,首先应通过数据清洗去除数据集中的噪声、重复和缺失值等不必要的数据,以保证数据的质量和准确性。具体来说,数据清洗包括以下几方面。

(1) 去重:去除选冶数据集中重复的数据行或数据列,以避免重复数据对后续选冶模型的分析造成影响,可以通过相似性度量的方法来解决选冶数据重复的问题。

相似性度量是对一个聚类结果中两个选冶数据点之间的相似性进行度量,度量方式有两种:用选冶数据点之间的距离来表示的相异度和用选冶数据点之间的相关性来表示的相似性。常用的相似性度量方法有:欧氏距离、曼哈顿距离等计算距离度量类方法,余弦相似度、相关系数法等相似度度量法。目前较为普遍使用的是基于欧几里得距离的去噪方法,该方法通过设置一个欧氏距离阈值来判断某个数据是否属于噪声数据。欧几里得距离可以通过以下等式计算:

$$d(x,v_i)=\sqrt{(x-v_i)^T(x+v_i)}$$

(5.1)

余弦相似度也是一种常见的相似度度量方法,这种方法利用两个样本之间形成的余弦值作为度量相似度的尺度,因此余弦相似度更加关注方向上的差异,其计算公式如下:

$$\text{sim}(X, Y) = \cos\theta = \frac{\mathbf{x} \cdot \mathbf{y}}{\|\mathbf{x}\| \|\mathbf{y}\|} = \frac{\sum_{i=1}^n x_i y_i}{\sqrt{\sum_{i=1}^n x_i^2} \sqrt{\sum_{i=1}^n y_i^2}} \quad (5.2)$$

余弦相似度的取值范围是 $[-1, 1]$ ，由余弦值的定义可知，余弦值越大，它们之间的夹角就越小，则这两个样本在这个方向上就越相似，反之则相反，这就是“余弦相似性”。

(2) 噪声过滤：通过设置阈值或使用机器学习算法等方法，可以去除采集到的选冶数据中的噪声或异常值。目前常用的智能选冶去噪算法大概可以分为五种：基于统计、基于邻近性、基于可视化、基于分类以及基于聚类。

基于统计的去噪方法通过建立概率模型来判断噪声点，采用不一致性检验的方法，假设正常的选冶数据处于模型中的高密度区域，噪声点则处于模型中的低密度区域。基于邻近性的去噪方法利用距离来判断噪声点，根据某个距离计算函数，计算选冶数据点之间的距离从而确定噪声值，也可以通过局部密度来探测异常值。基于可视化的方法利用图形图像，通过计算机图形技术、人机交互等技术，将传感器收集到的选冶数据映射为图像，直观地观察噪声点与正常数据之间的差别。基于分类的去噪方法通过构造模型来识别噪声点，通过训练集中的数据训练得到分类模型，用该模型对传感器收集到的数据进行分类，分类结束后不在任何类别中或具有异常的类属性的数据被判断为异常值。基于聚类的去噪算法可以同时进行分类与异常值检测的操作，在数据集大小方面的操作性较好，且时间复杂度与数据集的大小呈线性关系，算法较为高效。

(3) 缺失值处理：指对缺失值进行填充或删除，以提高选冶数据的完整性和可用性。对于缺失值处理的主体思路是填充缺失的数据，可以使用前后数据的均值补全时间序列的缺失值。若缺失值过多，一般使用平滑等处理，如滑动窗口。也可以依据统计学原理设置默认值，在数据缺失时用默认值填充。若数据无法补全则剔除数据，但仍保留不删除。

总体上说，对于选冶数据缺失的处理方法有三种，即删除、填充和不处理。数据填充技术大体上可以分为两种，一种方法是选冶数据操作人员或选冶领域的专家根据已有的知识经验人工地为缺失的数据填充合理的预期值；另一种方法是根据统计学原理，根据现有选冶数据的分布进行数据填充，如填充某个默认值或平均值。

为了提升数据的一致性，接下来对传感器收集到的选冶数据进行集成，主要通过建立数据体系、统一维度单位等途径来解决度量单位不一致的问题。数据维度不一致则可以采用降维方法，如主成分分析法、随机森林等；或者采用抽象方法，如汇总、离散化等方法解决；若选冶数据差异大则通过 Min-Max、z-score 等途径进行归一化。后续数据转换是将原始数据转换为更适合于机器学习算法的形式，这一步会将原始数据中选冶需要的特征缩放到同一尺度上以避免不同特征之间的尺度差异对模型的影响。接下来会具体介绍特征提取过程，最后就可以将数据导入选冶模型中运行和进行决策。

5.2.3 智能选冶数据特征提取

智能选冶数据特征提取的过程通常开始于数据采集，然后逐步过渡到数据预处理、特征选择、特征构造和特征转换等环节。在特征选择阶段，依据对选冶过程的专业知识和矿物冶炼的实践经验，挑选出那些对最终产品质量或者生产效率有显著影响的特征。通过相关性

分析或其他统计方法,可以识别并剔除冗余或无关紧要的特征,以降低模型复杂度。紧接着的是特征构造环节,这里可能需要基于领域内的具体知识来创造新的特征,以便更好地反映和解释数据中的内在规律。例如,可以结合多个参数构建出一个复合指标,或者根据选冶矿物的物理定律和化学反应原理推导出新的特征表达形式。最后是特征转换,这里的工作重点是将特征缩放到同一尺度上,避免不同特征之间的尺度差异对模型的影响。

完成特征提取后,所得到的特征即可输入各种机器学习算法或深度学习模型中,用于建立预测或分类模型,实现对选矿和冶炼过程的智能优化。这整个过程是迭代和反复的,随着新数据的不断涌入,可能需要回到前面的某个步骤进行调整和优化,以确保特征提取的持续准确性和有效性。智能选冶数据特征提取是提升选冶生产效率和质量的关键,未来随着人工智能技术的不断发展,特征提取方法将变得更加智能化和自动化,并进一步推动选冶向智能化发展。

5.3 基于信息融合的选冶破碎机、皮带机故障诊断

5.3.1 基于信息融合的选冶破碎机故障诊断

1. 破碎机及其损坏原理

选冶通常包括两个主要步骤:选矿和冶炼。选矿旨在将矿石中的有用矿物与杂质分离,以提高矿石的品位,减少后续冶炼过程中的能源和成本消耗。破碎机能够将从矿山采集的巨大岩石或矿石进行粗碎和细碎,将其破碎成符合磨矿或选矿要求的颗粒度,为后续的磨矿、浮选和冶炼等工艺过程提供合适的原料。破碎机可以适应各种矿石的硬度和黏附性,为后续的选择和冶炼提供了多样化的物料处理方式。通过选择不同类型和参数的破碎机,可以满足不同硬度和黏附性的矿石破碎要求,以确保选矿生产线的连续稳定运行。

破碎机根据不同的破碎原理和结构形式可以分为多种类型,如颚式破碎机、锤式破碎机、圆锥破碎机等。颚式破碎机,是采用间断破碎方式的代表,其机器结构相对简单,且机型较小。图 5-3(a)所示为颚式破碎机的常见结构,其工作原理是利用动颚和静颚的相对运动,将物料进行挤压破碎。锤式破碎机是靠篦板孔出料的,图 5-3(b)所示为锤式破碎机的常见结构。圆锥破碎机的工作原理是利用圆锥形的破碎腔室内壁和圆锥轴的相对运动,对物料进行挤压破碎,圆锥破碎机常用于破碎各种矿石和矿石中的原料。如图 5-3(c)所示,物料进入机器的进料口后,圆锥破碎机的圆锥体旋转并压缩物料,将物料破碎成所需粒度的颗粒。

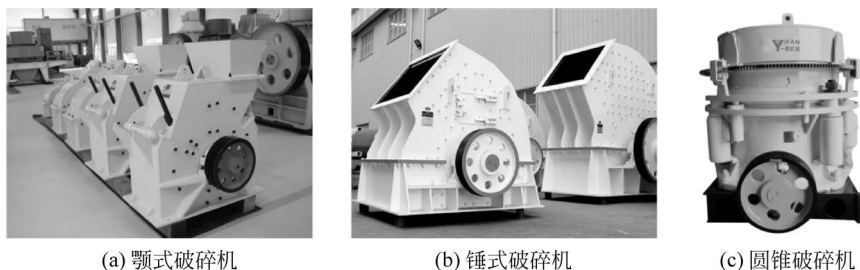


图 5-3 不同种类破碎机示意图

然而,破碎机在使用过程中可能会遭受各种损坏,其中包括零部件磨损、过载工作导致

的损坏等。破碎机损坏的原因涉及多方面,主要包括以下几个方面。油温过高:在破碎过程中,机器会产生大量的热量,如果油温过高,会导致油液黏度下降,摩擦阻力增加,进而影响偏心套的正常运转。油压异常:油压过低或过高都会对偏心套产生不利影响,油压过低可能会导致润滑不足,加剧磨损;而油压过高则可能会导致密封失效,引起内泄或外泄。轴承损坏:轴承是偏心套的关键部件之一,如果轴承出现损坏或磨损,将直接影响偏心套的运转,图 5-4 所示为破碎机损坏的轴承。物料堵塞:在破碎过程中,如果物料堵塞,会导致偏心套运转不畅,产生额外的压力和摩擦力,进而引发故障。

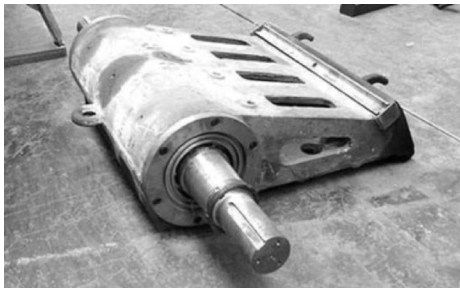


图 5-4 破碎机损坏的轴承

2. 破碎机故障检测方法

破碎机是选冶过程中常用的设备,故障的及时检测对于维护设备、保障生产效率至关重要。破碎机的故障检测方法多样,包括视觉检测、声音检测、振动检测、温度检测、润滑油检测等。

视觉检测:通过肉眼观察破碎机的外部情况,检查是否有异常磨损、变形、裂纹等现象,特别关注磨损部位、连接部位、传动部位、密封部位等,以判断设备是否存在潜在的故障隐患。声音检测:通过听觉对破碎机运行时的声音进行判断,异常的噪声可能来自轴承故障、齿轮间隙过大、传动带松动等问题。振动检测:利用振动传感器对设备的振动情况进行监测,异常的振动可能说明设备存在不平衡、松动、磨损等问题,从而及时发现并处理这些故障。温度检测:使用红外线测温仪或红外相机等工具检测破碎机各部位的温度,异常升高的温度可能表明轴承、电机等部件存在问题。润滑油检测:定期对设备的润滑油进行化验分析,以检测金属颗粒、水分等杂质,判断设备是否存在磨损或密封不良问题。

除了这些传统的方法外,还可以借助先进的传感器技术和数据融合监测系统进行故障检测。利用监测设备的运行参数、振动、温度、电流、压力等信息,结合数据融合分析和人工智能算法,能够准确地识别故障特征并提供预警提示,实现对潜在故障的早期发现和处理。

3. 基于数据融合的选冶破碎机故障诊断概论

基于数据融合的选冶破碎机故障诊断可以将来自不同传感器和监测设备的多源数据整合起来,从而全面、准确地分析和判断破碎机的运行状态,保障设备的稳定运行和安全生产。这种方法可以提高故障诊断的准确性和及时性,降低生产过程中的停机损失,增加设备维护的可预测性,从而提高生产效率和生产安全性。

首先,基于数据融合的选冶破碎机故障诊断可以实现对设备状态的全面监测。不同的传感器和监测设备可以收集到破碎机运行过程中的压力、声音、振动、温度等多种数据,这些数据反映了设备运行的全方位情况。图 5-5 为各种传感器实例图,通过对这些传感器数据进行融合分析,能够更全面地了解破碎机的工作状态,包括机械磨损情况、轴承状态、润滑油状况等,有助于准确诊断设备的健康状况。

其次,基于数据融合的选冶破碎机故障诊断也提高了故障诊断的及时性,因为传感器可以提供实时数据监测,一旦出现异常就可以发出预警,及时通知相关人员对设备进行检修和维护,避免故障进一步恶化,从而降低了生产过程中的停机损失。

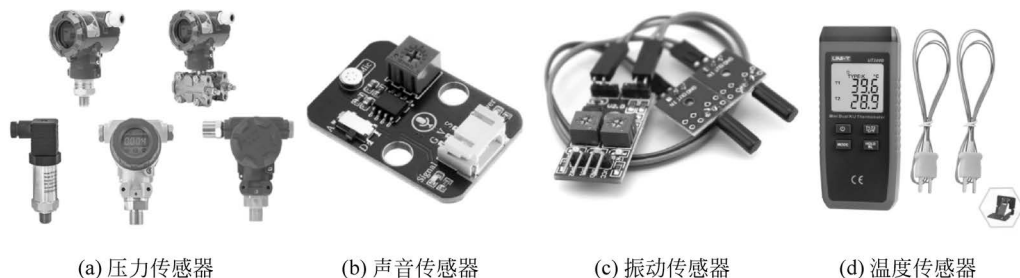


图 5-5 各种传感器实例图

4. 基于数据融合的破碎机故障诊断方法

基于数据融合的选冶破碎机故障诊断是一个高度综合性的技术领域,它结合了传感器技术、信号处理、机器学习和人工智能等多个学科的知识。在破碎机等工业设备的健康监测与故障诊断中,通过声音、振动、温度和润滑油等不同类型的监测数据,可以实现对设备运行状态的全面理解和及时故障预警,主要有以下方法。

1) 统计学方法

(1) 时频分析: 指将时间相关信号转换到频率域进行分析,识别异常的频率成分。时频分析(Time-Frequency Analysis, TFA)是处理非平稳信号的强大工具,特别适用于从包含多种成分的复杂信号中提取信息,如故障振动信号。在破碎机等旋转机械的故障诊断中,时频分析能够揭示设备运行状态中的瞬变特征和多成分交互作用,为故障的早期发现和诊断提供重要信息。

基于时频分析的数据融合故障诊断流程通常包括以下几个步骤。

- 数据采集和处理。
- 时频分析: 将预处理后的信号输入时频分析模块,常用的时频分析方法包括短时傅里叶变换(STFT)、小波变换(WT)、Wigner-Ville 分布(WVD)以及 Choi-Williams 分布等。
- 特征提取: 从时频分布中提取有助于故障识别的特征,这些特征可能包括特定频率带的能量、瞬时频率、时间-频率脊线等。
- 数据融合: 结合来自不同时频分析以及其他类型数据(如温度、润滑油分析)的特征,使用数据融合技术得到综合特征集。
- 故障诊断: 将融合后的特征输入分类器进行故障检测和诊断。
- 结果解释与决策: 根据分类器的输出确定是否存在故障,以及故障的类型和程度,并据此制定维护策略。

时频分析可以深入了解非平稳和非线性信号,可以捕捉到瞬时事件和频率变化,这对于故障的早期检测至关重要。数据融合增加了诊断系统的信息量和冗余度,提高了诊断结果的可靠性和准确性。通过融合不同类型的数据(如振动、声音、温度等),可以从多个角度监测设备的健康状况,实现更全面的监控和诊断。

(2) 相关性分析: 指评估不同传感器信号之间的相关性,判断异常状况。相关性分析是统计学中用于评估两个或多个变量之间相互依赖关系的方法。在选冶破碎机故障诊断领域,基于相关性分析的信息融合技术可以有效地结合来自不同传感器和数据源的信息,提高故障检测的准确性和可靠性。

以下是利用相关性分析进行信息融合及故障诊断的基本步骤。

- 使用传感器收集数据：首先，从安装在破碎机上的各种传感器（如振动、声音、温度、压力等）收集多维运行数据，这些数据反映了设备在不同工况下的表现，以及可能出现的异常情况。
- 数据处理：对原始数据进行必要的预处理操作，包括滤波、去噪、同步化处理、采样率转换等，以确保数据的质量和后续分析的准确性。
- 相关性分析：计算不同传感器数据之间的相关系数或互相关系数，这可以通过皮尔逊相关系数、斯皮尔曼相关系数（见图 5-6）或其他非线性相关度量来完成。相关性分析揭示了不同数据源之间的线性或非线性依赖关系，有助于识别出由特定故障引起的信号变化模式。

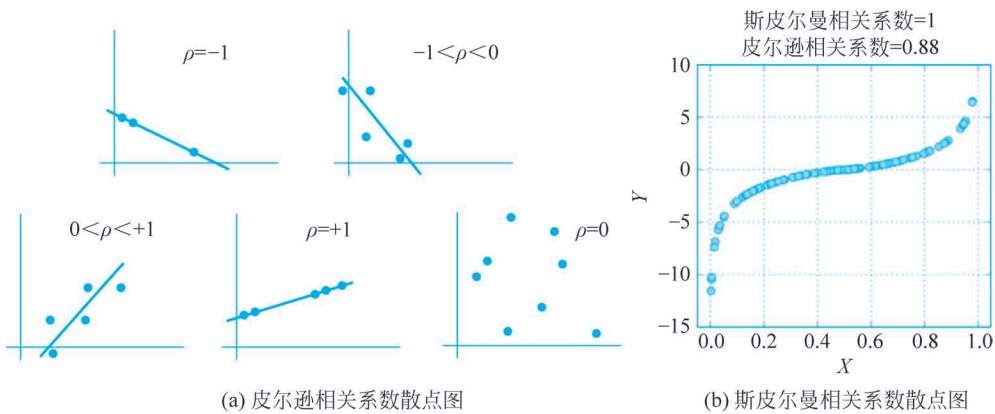


图 5-6 相关系数散点图

上述步骤完成了提取相关特征并进行训练和测试及决策制定等工作。此外，相关性分析还可以帮助确定维修的优先级和紧急程度。

相关性分析可以揭示不同数据源间的潜在联系，增加对设备状态的理解，从而提高故障检测的准确率。

2) 机器学习方法

(1) 卷积神经网络：卷积神经网络是一种深度学习模型，特别适合处理具有空间相关性的数据，如图像。在选冶破碎机故障诊断中，CNN 也可以被用来分析来自传感器的时序数据，尤其是当这些数据可以被转换成一维或二维的“图像”形式时，图 5-7 为卷积神经网络处理图像的过程示意图。

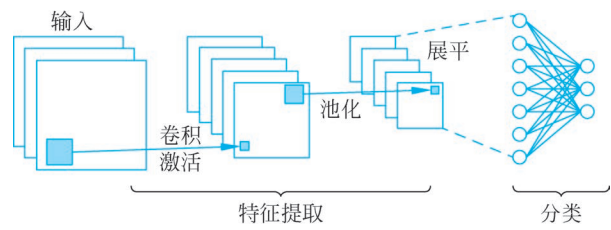


图 5-7 卷积神经网络处理图像的过程示意图

基于卷积神经网络的信息融合故障诊断流程包括以下步骤。

- 数据构造：从安装在破碎机上的多种传感器收集多源数据。

- 数据预处理：对原始数据进行滤波、去噪、归一化和同步处理等操作以确保数据的质量和一致性。
- 数据转换：将传感器数据转换为 CNN 可以处理的形式,对于时序数据,这可能意味着将数据重构成图像状,如绘制振动信号的频谱图,或者将时间序列数据转换为距离矩阵。
- 信息融合：利用 CNN 的卷积层自动从转换后的数据中提取特征,这些特征通常是多级别的,随着网络层次的加深,提取的特征也更加抽象和复杂。
- 模型训练：使用标注好的训练数据集来训练 CNN 模型,训练过程包括调整网络参数(权重和偏置)以最小化损失函数,通常使用反向传播和梯度下降算法。
- 检测：将新的测试数据通过同样的预处理和数据转换流程,然后输入已训练的 CNN 模型中进行故障检测和分类。
- 结果评估与决策：根据 CNN 模型的输出确定设备是否存在故障以及故障的类型和程度,这些信息用于指导维护决策和制定相应的维修策略。

(2) 循环神经网络(Recurrent Neural Network, RNN)：擅长处理序列数据,如时间序列分析。循环神经网络是一种专门用于处理序列数据的神经网络结构。在工业故障诊断领域, RNN 可以有效地处理时间序列数据,从而对设备的工作状态进行实时监控和预测。

基于循环神经网络的信息融合选冶破碎机故障诊断方法可以分为以下几个步骤。

- 数据收集与预处理：从选冶破碎机的传感器中收集原始数据,如振动信号、温度信号等。对数据进行预处理,包括滤波、降噪、归一化等操作,以便于后续的特征提取和模型训练。
- 提取特征：从预处理后的数据中提取有助于故障诊断的特征。这可以通过时域分析(如统计特征)、频域分析(如傅里叶变换)或时频域分析(如小波变换)等方法实现。
- 构建循环神经网络模型：设计一个适用于故障诊断任务的 RNN 结构,常见的 RNN 结构有长短时记忆网络(Long Short-Term Memory, LSTM)和门控循环单元(Gated Recurrent Unit, GRU)。这些结构可以捕捉时间序列数据中的长期依赖关系。
- 融合传感器数据：将不同传感器收集到的数据进行融合,以获得更全面的设备状态信息,这可以通过多通道 RNN、注意力机制或其他信息融合方法实现。
- 模型训练与验证：使用带有标签的数据集对 RNN 模型进行训练,在训练过程中,通过调整模型参数以最小化损失函数,使模型能够准确地识别不同类型的故障。同时,使用验证集对模型的性能进行评估。
- 计算概率：将训练好的 RNN 模型应用于实际的选冶破碎机故障诊断场景,根据实时监测数据,模型可以预测设备的故障类型和发生概率,从而实现对设备故障的及时预警和维护。

(3) 支持向量机(Support Vector Machine, SVM)：是一种用于分类和回归分析的监督学习模型。SVM 的基本原理是在一个高维空间中寻找一个超平面,该超平面可以将不同类别的数据点有效地分隔开来。SVM 分类原理图如图 5-8 所示。图 5-8 中,三角形和圆形表

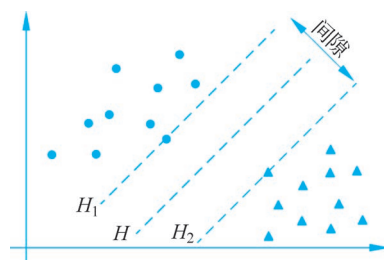


图 5-8 SVM 分类原理图

示不同类样本；平面 H 表示最优决策超平面，能够将两种样本类型进行最佳分类；平面 H_1 和 H_2 表示支撑平面。SVM 算法首先是计算所有三角形和圆形到 H 的距离，然后找出距离 H 最近的样本点，接着在最近样本点中找到距离最大的样本点，使支持向量到 H 的最小距离最大，这样就可以将两类数据完全分开了。假设两类样本为： $(x_i, y_i), i=1, 2, \dots, n, x_i \in R^d, y_i \in \{+1, -1\}$ ，最优决策超平面 H 为： $\omega x + b = 0$ ，其中

$\omega \in R^d$ ，通常与最优决策超平面 H 正交。构造两个与 H 平行的支撑平面，即

$$H_1: \omega^T x + b = +1 \tag{5.3}$$

$$H_2: \omega^T x + b = -1 \tag{5.4}$$

当 H 将两类样本数据分开时，数据边界样本到 H 距离最小的点就是支持向量，支持向量是支撑平面上的点。计算支持向量到最优决策超平面 H 的距离长度：

$$d = \frac{\omega_0^T x + b_0}{\omega_0^T} = \frac{1}{\omega_0} \tag{5.5}$$

式中， ω_0 是超平面 H 的权向量； b_0 是超平面 H 的偏置。 $d_{\min}$ 为 $1/\|\omega\|$ ，而两类样本数据之间的距离长度则是两倍的 $d_{\min}$ ，表示为 $2/\|\omega\|$ ，若要完全准确地将两类样本进行分类，就应该求得最大化后的 $1/\|\omega\|$ ，即

$$\min_{\omega, b} \frac{1}{2} \omega^2 \tag{5.6}$$

$$\text{s. t. } y_i(\omega^T x + b) \geq +1, \quad i=1, 2, \dots, n \tag{5.7}$$

此时把求解问题转换成了一个凸优化问题，再引入 Lagrange 函数将其转换成对偶优化问题，即影响参数只含有 Lagrange 因子，其内积形式方便数据被核函数替换，得到最优决策超平面 H 的函数关系：

$$f(x) = \text{sgn}\{\omega x + b\} = \text{sgn}\left\{\sum_i^n \alpha_i^* y_i(x_i, x_j) + b^*\right\} \tag{5.8}$$

实际问题中大多为非线性，需考虑采用非线性模型来求解，而 SVM 方法则是引入核函数 $\phi()$ 将低维空间的非线性问题映射到高维空间来构建线性分类模型，则最优分类面为

$$\omega \phi(x) + b = 0 \tag{5.9}$$

3) 实例分析(某矿山破碎机故障诊断)

(1) 破碎机的测点分布。

PXZ1200/160 破碎机由传动、机座、偏心套、破碎圆锥、中架体、横梁、原动机、油缸、液压系统、滑车、电气系统和干、稀油润滑等部分组成。该破碎机的具体参数如表 5-2 所示。

表 5-2 某破碎机的具体参数

规格型号	给料口/mm	排料口/mm	排料口范围/mm	生产率/(t/h)	电机功率/kW	重量/t
PXZ1200/160	1200	160	160~195	1050~1195 1250~1480	260 310	229

(2) 破碎机故障监测。

破碎机通过传感器完成数据采集,传感器测点分布说明如表 5-3 所示,包括偏心套故障(H_1)、轴承磨损(H_2)、平行轴缺油(H_3)等故障类型。

表 5-3 传感器测点分布说明

测 量 点	传感器类型	测 量 参 数
偏心套	振动传感器	振动频率
轴承	振动传感器	振动频率
动锥	振动传感器	振动频率
破碎机底座	声音传感器	振动频率
润滑油油箱	温度传感器	润滑油油温
回油孔处	温度传感器	回油油温
冷却水管出口处	温度传感器	冷却管出油油温
电机轴承处	温度传感器	轴承温度
润滑油油箱	压力传感器	润滑油油压
液压缸	压力传感器	液压缸油压
过滤器	压力传感器	过滤前后压差

温度特征数据集如下,其训练样本空间数据如表 5-4 所示。

$$I_3 = \{C_1, C_2, C_3, C_4\} \tag{5.10}$$

其中, C_1 为润滑油油温, C_2 为冷却管出油油温, C_3 为回油油温, C_4 为轴承温度。

表 5-4 温度特征训练样本空间数据

故 障 类 型	C_1	C_2	C_3	C_4
H_1	0.7742	0.3636	0.6897	0.5455
H_1	0.9355	0.3636	0.4138	0.2727
H_1	0.7742	0.2727	0.6552	0.5455
H_1	0.8387	0.3939	0.5172	0.3636
H_2	0.7419	0.3939	0.3448	0.8182
H_2	1.0000	0.3333	0.3793	0.3182
H_2	0.9032	0.3939	0.3448	0.3636
H_2	0.7419	0.3333	0.3103	0.6818
H_3	0.0645	0.0303	0.0000	0.1364
H_3	0.0000	0.1818	0.0345	0.0909
H_3	0.0000	0.0000	0.0000	0.2273
H_3	0.0000	0.1515	0.0345	0.2273

压力特征数据集如下,其训练样本空间数据如表 5-5 所示。

$$I_4 = \{D_1, D_2, D_3, D_4\} \tag{5.11}$$

其中, D_1 为润滑油油箱油压, D_2 为液压缸油压, D_3 为过滤前油压, D_4 为过滤后油压。

表 5-5 压力特征训练样本空间数据

故 障 类 型	D_1	D_2	D_3	D_4
H_1	0.2273	0.4091	0.1429	0.1538
H_1	0.1818	0.5000	0.0714	0.1154
H_1	0.4091	0.0909	0.0714	0.2231

续表

故障类型	D_1	D_2	D_3	D_4
H_1	0.0909	0.2273	0.2143	0.1538
H_2	0.2273	0.1818	0.7857	0.6154
H_2	0.2273	0.0455	0.7143	0.4615
H_2	0.1818	0.0909	0.5714	0.3846
H_2	0.2727	0.0000	0.5000	0.6154
H_3	0.6364	0.9091	0.5714	0.7692
H_3	0.7727	0.8636	0.7143	0.8462
H_3	0.9091	0.6364	0.2857	1.0000
H_3	1.0000	0.6364	0.5000	0.4615

振动特征数据集如下,其训练样本空间数据如表 5-6 所示。振动特征数据集 I_1 表达式如式(5.12)所示。

$$I_1 = \{A_1, A_2, A_3, A_4, A_5, A_6, A_7\} = \left[\frac{E_0}{E}, \frac{E_1}{E}, \frac{E_2}{E}, \frac{E_3}{E}, \frac{E_4}{E}, \frac{E_5}{E}, \frac{E_6}{E}, \frac{E_7}{E} \right] \quad (5.12)$$

其中, $A_i (i=0,1,\cdots,7)$ 为相应的特征 E_i 进行标准化后的振动特征值, $E = \sum_{i=0}^7 E_i, E_0 \sim E_7$ 分别表示振动幅值、振动频率、振动加速度、振动相位、振动波形,振动频谱、振动峰值因子、振动峭度。

表 5-6 振动特征训练样本空间数据

故障类型	A_1	A_2	A_3	A_4	A_5	A_6	A_7
H_1	0.0424	0.2990	0.1622	0.2411	0.0012	0.0016	0.1929
H_1	0.0408	0.3009	0.1590	0.2473	0.0014	0.0019	0.1707
H_1	0.0392	0.3028	0.1558	0.2535	0.0016	0.0022	0.1485
H_1	0.0404	0.3259	0.1490	0.2447	0.0014	0.0015	0.1800
H_2	0.0202	0.0110	0.4442	0.0920	0.0015	0.0015	0.3233
H_2	0.0343	0.0201	0.4486	0.1152	0.0015	0.0012	0.2958
H_2	0.0366	0.0125	0.4318	0.0878	0.0005	0.0014	0.3214
H_2	0.0326	0.0201	0.4445	0.1360	0.0015	0.0015	0.2622
H_3	0.3857	0.0710	0.3480	0.0863	0.0001	0.0006	0.0688
H_3	0.3922	0.0712	0.3121	0.1010	0.0002	0.0007	0.0815
H_3	0.4063	0.0789	0.3334	0.0843	0.0002	0.0006	0.0673
H_3	0.3432	0.0884	0.3737	0.0894	0.0005	0.0004	0.0728

声音特征数据集 I_2 表达式如下,其测试样本空间数据如表 5-7 所示。

$$I_2 = \{B_1, B_2, B_3, B_4, B_5, B_6, B_7\} = \left[\frac{E_0}{E}, \frac{E_1}{E}, \frac{E_2}{E}, \frac{E_3}{E}, \frac{E_4}{E}, \frac{E_5}{E}, \frac{E_6}{E}, \frac{E_7}{E} \right] \quad (5.13)$$

其中, $B_i (i=0,1,\cdots,7)$ 表示相应的特征 E_i 进行标准化后的声音特征值, $E = \sum_{i=0}^7 E_i, E_0 \sim E_7$ 分别表示声音振幅、声音频率、声音持续时间、声音间隔、音色、音调、声音能量、声音熵。

表 5-7 声音特征测试样本空间数据

B_1	B_2	B_3	B_4	B_5	B_6	B_7	诊断结果
0.0160	0.5227	0.0011	0.0020	0.0242	0.0062	0.0007	H_1
0.0305	0.0292	0.0085	0.3025	0.5290	0.0208	0.0102	—
0.5007	0.2189	0.0186	0.1078	0.0405	0.0203	0.0387	H_3

SVM 核函数及相关参数的选取：选择 RBF 函数作为 SVM 模型的核函数，即

$$K(\mathbf{x}, \mathbf{x}_i) = \exp\left\{-\frac{\|\mathbf{x} - \mathbf{x}_i\|^2}{\sigma^2}\right\} \quad (5.14)$$

最终构造的 SVM 模型为

$$f(\mathbf{x}, \alpha) = \text{sgn}\left\{\sum_{i=1}^s \alpha_i \exp\left\{-\frac{\|\mathbf{x} - \mathbf{x}_i\|^2}{\sigma^2}\right\} - b\right\} \quad (5.15)$$

SVM 惩罚因子 $C=8$ ，核函数参数 $\sigma=0.078$ 。

选用了“一对一”多分类模型进行破碎机故障诊断研究。故障诊断识别框架 $\Theta = \{H_1, H_2, H_3, H_4, H_5, H_6, H_7\}$ ，故障类型包括偏心套故障(H_1)、轴承磨损(H_2)、平行轴缺油(H_3)，建立 SVM 二分类器。

根据各特征数据中的测试样本得出 SVM 初步诊断结果，构建出特征空间证据体，从而计算各证据体的基本可信度 M_1 ，由此可以得出，故障诊断正确率明显提高，融合效果明显。其中计算的第 6 组数据(针对某一特定测试样本所计算得到)的信任度区间和基本可信度分配的一组具体数值信任度区间为 $[0.1533, 0.3168]$ 、 $[0.5818, 0.7453]$ 、 $[0.4233, 0.5868]$ 、 $[0.0202, 0.1837]$ 、 $[0.0649, 0.2284]$ 、 $[0.0565, 0.2200]$ 、 $[0.2481, 0.4116]$ ， $m(\Theta)=0.1635$ ，最大信任度为 $m(H_w)=0.5818$ ，故 H_2 为故障目标类。 $m(H_w)-m(H_i)=0.5818-0.4233=0.1585<\epsilon_1=0.2, i \neq 2$ ，故不符合规则。 $m(H_w)-m(\Theta)=0.5818-0.1635=0.4183>\epsilon_2=0.4$ ，故符合规则。 $m(\Theta)=0.1635, \epsilon_3=0.2$ ，故符合规则。

通过对实际选冶破碎机进行信号获取、特征提取，可以构成数据样本集，将实际信号数据输入 SVM 模型就可以对破碎机故障类型进行识别诊断，为预警系统提供基础，减少维修成本及故障发生率，降低破碎机故障分类诊断的错误率，提高其分类结果的有效性和可靠性，有效提高矿山企业设备管理效率。

5.3.2 基于信息融合的选冶皮带机故障诊断

1. 皮带机结构及其故障损坏原理

作为原矿输送的主要工具，皮带机高效的物料输送功能确保了选冶厂始终有稳定的原材料供应，皮带机的运行直接影响选冶生产的效率。皮带机的工作原理基于输送带的连续运动，其核心组成部分包括输送带、驱动系统、滚筒和托辊、调节装置以及清洁装置。皮带机的大致结构如图 5-9 所示。在工作时，通过电动机等驱动装置带动输送带的运动，输送带上的物料随之被带动向前输送。滚筒和托辊支撑和引导输送带的运动，调节装置用于维持输送带的张紧度，而清洁装置则有助于防止物料残留和减少设备磨损。

不同类型的皮带机虽然在结构和参数等方面存在差异，但一般均包括输送带、减速器、滚筒、托辊等关键部件。在皮带机运行过程中，受长运距、大运量、高运速、恶劣运行环境、部件分布广且耦联关系复杂等因素影响，其关键部件易产生多种故障或损伤，如输送带撕裂与



图 5-9 皮带机的大致结构

打滑、托辊磨损与卡死、滚筒磨损与破裂、齿轮箱磨损与断裂等,严重影响皮带机的健康状况,威胁其安全可靠运行。

以下为输送带、减速器、滚筒、托辊等皮带机关键部件所存在的故障种类、原因及危害分析。

(1) 输送带。输送带的故障主要包括输送带打滑、跑偏、损伤和堆料撒料等。输送带打滑会导致输送效率下降、物料堆积等问题,物料堆积又会造成输送带表面不平整或堵塞,进一步加剧打滑的程度,长时间打滑则会加剧输送带与托辊、滚筒的过度摩擦,引起温度上升,进而引发潜在火灾的发生。输送带跑偏也会导致物料堆积,并会加剧输送带的磨损。堆积物的堵塞还可能导致输送带破损或断裂,增加了设备维护和更换的成本。

(2) 减速器。减速器齿轮箱的故障一般发生在其内部齿轮、轴承等部件上,齿轮和轴承的错误安装、润滑不良等原因,会加剧自身的磨损。除皮带机过载、过热外,齿轮的磨损和轴承的故障也均有可能演变为齿轮的断裂,这将导致皮带机无法工作,造成经济损失。齿轮的断裂也可能导致输送带突然停止或撕裂,进而引发多个部件的故障。

(3) 滚筒。滚筒的故障包括筒体开焊、包胶破损、包胶脱落等。滚筒作为皮带机主要的传动部件,长期工作在复杂而恶劣的工况环境中,由于负荷大、负载不平衡等因素影响,经常产生包胶破损、脱落等故障,易形成安全隐患,酝酿安全事故。

(4) 托辊。托辊的故障包括磨损和卡死。托辊磨损或卡死一方面会导致输送带不平整或不稳定,使得输送带在运行过程中产生跳动或偏移,进而可能引发物料堆积、洒落或堵塞等问题;另一方面会加剧输送带的磨损,导致输送带撕裂或断带。

2. 皮带机故障诊断中常见的信息融合技术

由于选冶皮带机所处工况环境恶劣,采集到的原始信号一般含有大量噪声,为此需要使用数据处理方法提取、变换或融合信号中有用的信息。从检测融合的角度出发,信息融合包括四种结构模型:并行结构、串行结构、分散式结构和树状结构。并行结构是通过各个传感器的原始数据采集、局部判断之后通过融合中心实现全局判断的一种检测融合结构。串行结构是下游传感器通过上游传感器局部诊断结果和自身原始数据进行融合判断,向下递推,直至得到最终融合结果的检测融合结构。分散式结构是各个传感器分别对自身原始数据进行最终判断的检测融合结构。树状结构是并行结构和串行结构相结合的混合结构,由树枝开始依次向下,在各个节点处分别实现检测融合,直至树根节点实现全局检测融合过程。

对于皮带故障检测而言,检测融合最常采用的是并行结构和分散式结构,如图 5-10 所示。其中,并行结构用于多传感器判断单个事故,而分散式结构用于单传感器判断单个事故。

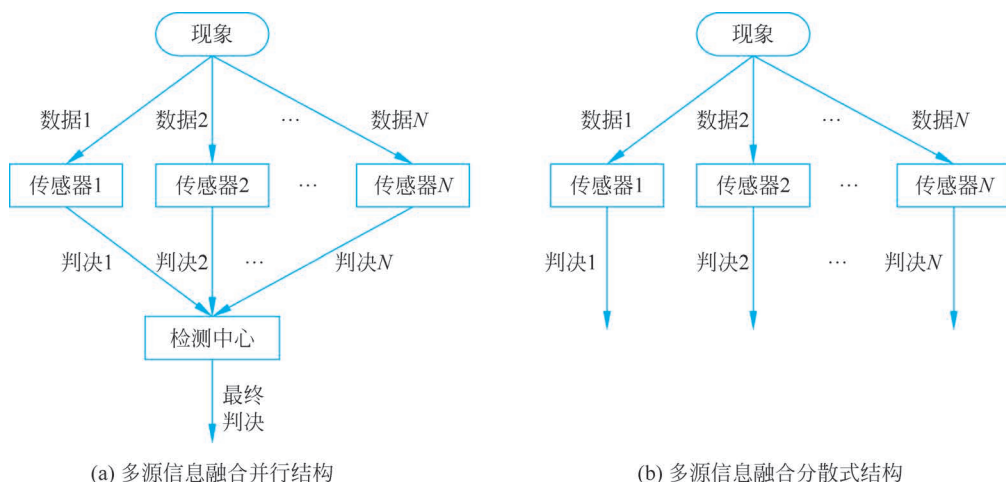


图 5-10 检测融合方式结构图

目前,在并行的多传感器信息融合方法中,常用的方法有加权平均法、神经网络、贝叶斯推理和 D-S(Dempster-Shafer Method)证据理论等。其中,神经网络需要适当的训练规则,且针对复杂情况下的泛化能力较弱。贝叶斯推理则需要先验概率的参与,且只能判断皮带机是否发生故障。D-S 证据理论克服了贝叶斯推理先验概率不易获得的缺点,是对贝叶斯推理的扩充,考虑了条件概率的不确定度,给每个故障的前提条件都设置了可信度。综合来看,D-S 证据理论适用于皮带机故障检测的多传感器并行结构下的多源信息融合。

3. 实例分析(基于 D-S 证据理论的选冶皮带机故障检测系统)

1) 理论方法

D-S 证据理论是由 Dempster 提出,在其学生进一步研究下形成 Dempster 合成法则。证据理论深化了命题和集合之间的关系,其组合规则为:假设识别框 M 下的两个证据为 P_1 和 P_2 ,其相对应的信任度函数值分配表示为 m_1 和 m_2 ; 焦元分别为 J_1 和 J_2 ,则通过式(5.17)定义函数 $m: 2M \rightarrow [0,1]$,表示其融合后的信任度函数分配为

$$m(A) = \begin{cases} 0, A = \emptyset \\ \frac{1}{1-K} \sum_{J_1 \cap J_2} m_1(J_1) \cdot m_2(J_2), A \neq \emptyset \\ K = \sum_{J_1 \cap J_2} m_1(J_1) \cdot m_2(J_2) < 1 \end{cases} \quad (5.16)$$

式中, K 为不确定因子,反映证据的冲突程度; $\emptyset$ 为空集; $1/1-K$ 为归一化因子; $m(A)$ 为对应 m 的证据基本信任度函数值。根据式(5.16)分别计算出各证据之间的基本可信度,并且利用其组合规则求取所有证据的信任度函数之后,根据设定的决策规则,选取概率最大的目标。

在皮带机设备的故障诊断过程中,需要诊断不同故障或是多种故障同时发生的征兆,这就需要对融合信息进行故障诊断。将多种故障征兆信息融合之后求取最大的信任度函数

值,并根据输入的判断规则,判断皮带机的故障类型。首先根据 D-S 证据理论建立故障识别框架,并且采集可能发生的故障症状特征,构建故障识别数据库。首先假设输送机中共有 n 个智能传感器、 f 个故障类型,通过 D-S 证据理论构建信任度函数分布矩阵 $\mathbf{S}$,如式(5.17)所示:

$$\mathbf{S} = \begin{bmatrix} \mathbf{S}_1 \\ \mathbf{S}_2 \\ \vdots \\ \mathbf{S}_n \end{bmatrix} = \begin{bmatrix} s_{11} & s_{12} & \cdots & s_{1f} \\ s_{21} & s_{22} & \cdots & s_{2f} \\ \vdots & \vdots & \ddots & \vdots \\ s_{n1} & s_{n2} & \cdots & s_{nf} \end{bmatrix} \quad (5.17)$$

矩阵 $\mathbf{S}$ 中的每一个证据元素 s_{nf} 表示的是,第 n 个传感器可能发生第 f 种故障的信任度分配函数,则得到同一个传感器计算的信任度分配函数的总和为 1,用式(5.18)表示为

$$\begin{cases} s_{n1} + s_{n2} + s_{n3} + \cdots + s_{nf} = 1 \\ n = 1, 2, \dots, f \end{cases} \quad (5.18)$$

式中, s_{n1} 为第 n 个传感器中发生的第 1 种故障的信任度分配函数,其余以此类推。并且,将其中某一行的转置向量乘以另一行,可以得到一个新的矩阵 $\mathbf{Y}$,如式(5.19)所示:

$$\mathbf{Y} = \mathbf{S}_n^T \mathbf{S}_f = \begin{bmatrix} s_{n1}S_{j1} & s_{n1}S_{j2} & \cdots & s_{n1}S_{jn} \\ s_{n2}S_{j1} & s_{n2}S_{j2} & \cdots & s_{n2}S_{jn} \\ \vdots & \vdots & \ddots & \vdots \\ s_{nj}S_{j1} & s_{nj}S_{j2} & \cdots & s_{nj}S_{jn} \end{bmatrix} \quad (5.19)$$

式中, $\mathbf{S}_n^T$ 为转置向量。

根据 D-S 证据理论的组合规则,采集 3 个及 3 个以上的故障类型特征时,需要先选取 2 个故障类型特征进行组合,然后将组合的结果与第 3 个故障类型特征进行重新组合,根据采集的故障数据以此类推,将一个新的组合结果与剩下的故障类型特征证据进行组合,直到最后得出结果。根据相应的判断规则,选取融合之后发生概率最大的为皮带机的主要故障类型。

2) 仿真实验

利用 MATLAB 平台进行 D-S 证据理论实现仿真实验,通过模拟仿真传感器采集的故障特征证据融合,得出输出融合后的可信度,如表 5-8 所示。

表 5-8 两次融合后的可信度

结果故障	f_1	f_2	f_3
$S_1 \oplus S_2$	0.79	0.12	0.08
$S_1 \oplus S_2 \oplus S_3$	0.981	0.031	0.007

从仿真实验结果可以看出,第二次融合的结果相较第一次来说有着更高的准确性。从证据特征融合结果来看,皮带机故障类型为 f_1 的可能性更高,更加趋近于 1,而其他故障类型的信任度函数值在逐渐减少。这与 D-S 证据理论的算法一致,随着证据融合次数增加,不确定性减少,其信任度函数值会向判断类型靠拢。从表 5-8 中的数据可以看出,基于 D-S 证据理论的皮带机故障智能检测系统的准确性和实用性较高,可以有效预警皮带机运行过程中出现的故障,保障煤矿企业的效益,确保系统安全可靠运行。

5.4 基于信息融合的冷风机与温度场工况分析

5.4.1 基于信息融合的冷风机工况分析

1. 冷风机概述

冷风机的主要功能包括提供足够的通风量,稀释并排除有害气体(如甲烷、一氧化碳等),控制选冶工作环境的温度和湿度。冷风机通常分为轴流式和离心式两种类型,图 5-11 为常见的冷风机示意图,其中轴流冷风机更为常见,它能够提供大流量的气流,适应选冶工作环境的通风需求。根据选冶的具体需要,冷风机可能配置有扩散器、消声器等附件,以优化通风性能和降低噪声。现代冷风机集成了智能化控制系统,能够实时监测风速、风压、温度等参数,并根据矿井内部情况自动调节运行状态。

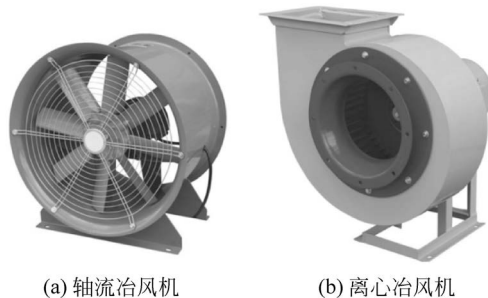


图 5-11 常见冷风机示意图

冷风机的工作原理主要是通过叶轮的旋转运动产生气流,并通过出口排出,同时形成负压吸入新的空气,实现连续的气流传递。具体来看,风机的核心部件是叶轮,它通过电动机或发动机的驱动产生旋转运动。叶轮的设计通常包括一系列叶片,这些叶片在旋转时能够推动空气移动。当叶轮旋转时,它会推动周围的空气,形成气流。这个过程是通过叶片与空气的相互作用实现的,叶片的形状和角度都会影响气流的速度和方向。产生的气流会通过风机的出口被排出,这个过程中,由于气流的流出,会在风机内部形成负压区域。为了平衡压力差,新的空气会被吸入风机,形成一个连续的气流循环。

2. 冷风机损坏机理

冷风机损坏机理通常指的是导致冷风机性能下降或完全失效的各种原因和过程,常见的冷风机损坏机理包括以下内容。

(1) 磨损: 长时间运转的风机,其内部组件特别是叶片和叶轮会因为摩擦、冲击和粒子冲刷而逐渐磨损。

(2) 腐蚀: 风机材料可能受到腐蚀性气体或水蒸气的侵蚀,导致叶轮和其他关键部件的金属材质发生化学反应,从而减弱结构强度。

(3) 疲劳断裂: 风机在长期运行下,叶轮等旋转部件会经历周期性载荷的作用,长时间的交变应力可能导致材料疲劳,进而产生裂纹甚至断裂。

(4) 温度效应: 高温环境下工作可能导致部分风机材料的性能降低,例如塑料零件可能会变形,润滑油可能会失效,导致轴承损坏。

综上所述,冷风机的损坏机理是多方面的,涉及设计、材料、制造工艺、运行环境等多个因素。为了确保冷风机的可靠运行,需要对这些潜在的损坏机理进行深入分析,并采取相应的预防和维护措施。

3. 冷风机故障检测方法

冷风机是选冶工业应用中不可或缺的设备,对冷风机进行有效的故障检测至关重要,它

是确保系统持续、安全和高效运行的重要因素。常见的冷风机故障检测方法包括以下内容。

(1) 振动分析：通过测量和分析风机的振动信号，可以识别出不平衡、对中不准、轴承损坏等机械问题。使用振动传感器（如加速度计）来收集数据，并通过频谱分析来确定异常振动的来源。

(2) 温度监测：电机过热、摩擦增加等问题可以通过监测温度来发现。安装温度传感器（如热电偶、红外传感器）来监测关键部件的温度，并设置报警阈值。

(3) 声音分析：异常声音（如尖叫声、嘶嘶声或敲击声）可能是故障的早期迹象，图 5-12(a)所示为声学分析仪，实施定期听觉检查或使用声音分析仪器来检测异常。

(4) 无损检测：使用超声波、X 射线或磁粉等技术来检测内部缺陷或裂纹。特别适用于检查焊接接缝、铸件和叶片等部件的完整性，图 5-12(b)所示为无损检测仪。

(5) 数据分析和预测维护：使用数据采集和分析技术来预测潜在故障，图 5-12(c)所示为检测系统。结合历史数据和机器学习算法，建立预测模型来指导维护决策。



图 5-12 故障检测仪器

4. 基于数据融合的冷风机故障诊断概论

基于数据融合的冷风机故障诊断是一个集成多源信息以检测和识别系统中潜在问题的过程，这一过程充分利用了来自不同传感器和数据源的信息，实现了冷风机运行状态的综合评估。基于数据融合的冷风机故障诊断的主要步骤包括：①从冷风电机组的多个传感器和监测点收集原始数据，包括但不限于振动传感器、温度传感器、油液分析、电机电流和功率等。②故障检测与诊断：利用融合后的数据，采用机器学习、深度学习或其他智能算法进行故障检测与诊断。③决策支持：故障诊断结果为运维人员提供决策支持，指导维修或更换部件。基于数据融合的冷风机故障诊断的具体融合方法如下。

1) 传感器数据融合

在风机系统中部署多种类型的传感器（如振动传感器、温度传感器、压力传感器、流量传感器等），收集不同维度的数据。然后对这些数据进行同步和预处理，如滤波、去噪和归一化。接下来，通过特征提取技术从原始数据中获取有用的信息，例如从振动信号中提取幅值、频率、波形指标等特征。最后，使用数据融合算法（如卡尔曼滤波器、贝叶斯网络、D-S 证据理论等）将不同传感器提供的特征数据进行综合分析，以获得更准确的故障状态估计。

应用数据融合技术（如卡尔曼滤波器、贝叶斯网络、模糊逻辑等）可以整合不同传感器的特征数据，增强故障信号的信噪比，降低误报率和漏报率。结合数据融合的结果，采用更高级的算法（如神经网络）进行进一步分析，以精确定位故障类型和程度。根据高级诊断结果，制订维护或修复计划。在此过程中，可以参考专家系统或历史维护记录来辅助决策。维护

完成后,将实际维护情况反馈到系统中,不断更新和优化故障诊断模型,提升诊断精度。

2) 模型级数据融合

模型级数据融合可以结合不同的故障诊断模型,通过集成学习方法(如 Bagging、Boosting、Stacking 等)将它们组合成一个更强大的模型,以提高诊断性能,模型级数据融合机理如图 5-13 所示。

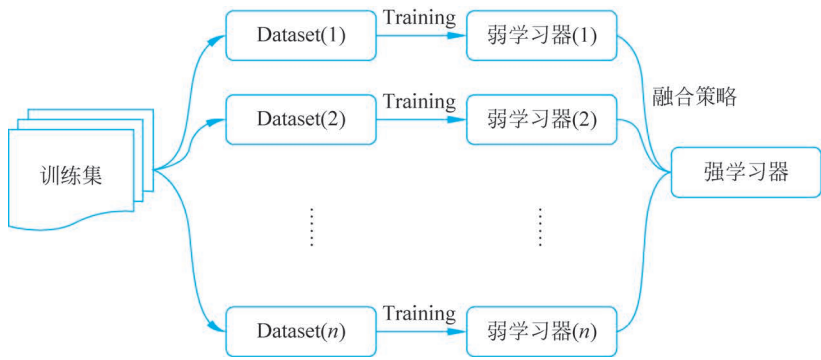


图 5-13 模型级数据融合机理

Bagging 是一种通过自助采样产生多个训练集,然后训练多个基学习器,并通过投票或平均来进行预测的方法。对于冶风机故障诊断,可以训练多个基于不同特征选择或采样方式的模型,然后综合它们的预测结果。Boosting 串行训练多个基学习器,每个学习器都会对前一个学习器的预测错误进行更多的关注,从而逐步提升模型的预测能力。对于冶风机故障诊断,可以使用各种提升算法(如 AdaBoost、Gradient Boosting 等)来结合多个弱学习器,从而构建更强大的集成模型。Stacking 是一种更为复杂的模型级数据融合方法,它通过将不同模型的预测结果作为新的特征输入,再训练一个元模型来做最终的预测。

3) 机器学习与统计学融合

应用机器学习算法(如支持向量机、随机森林、深度学习等)对历史数据进行训练,建立故障诊断模型,多源异构数据融合经典结构如图 5-14 所示。同时,使用统计学方法对数据的分布和趋势进行分析,为机器学习模型提供先验知识。

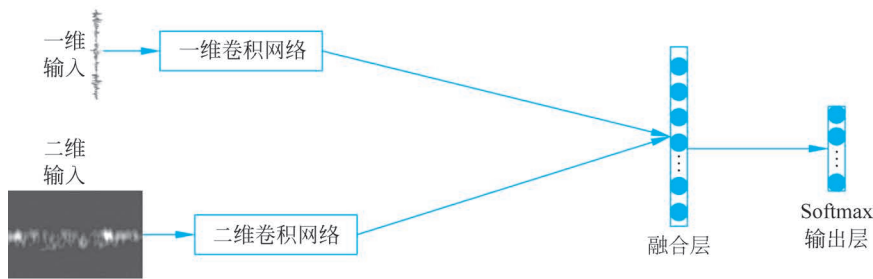


图 5-14 多源异构数据融合经典结构

总之,基于数据融合的冶风机故障诊断方法通过整合多种监测数据和分析技术,能够全面评估设备的健康状况,提前发现潜在问题,并指导维护决策。

4) 基于粗糙集理论的融合

论域与概念: 设 M 是研究对象组成的非空有限集合,称为论域。论域 M 的任一子集

$X \subseteq M$, 称为论域 M 的一个概念。给定 1 个论域 M 和 M 上的一簇等价关系 T , 称二元组 $K = (M, T)$ 是关于论域 M 的一个知识库。给定 1 个论域 M 和 M 上的一簇等价关系 T , 若 $P \subseteq T$, 且 $P \neq \emptyset$, 则 $\cap P$ 仍然是论域 M 上的一个等价关系, 称为 P 上的不可分辨关系。

在 RS 理论中, 称四元组 $KR = (M, A, V, f)$, 是一个关系数据表(知识表达系统), 其中, M 为论域, 即对象的非空有限集; A 为属性集, 即属性的非空有限集。

给定知识库 $K = (M, T)$, 则任意 $X \subseteq M$ 和论域 M 上的一个等价关系 $R \subseteq \text{IND}(K)$, 定义子集 X 关于知识 R 的下近似和上近似分别为

$$\bar{R}(X) = \{x \mid (\forall x \in M) \wedge ([x]_R \subseteq X)\} = \bigcup \{Y \mid (\forall Y \in M/R) \wedge (Y \subseteq X)\} \quad (5.20)$$

$$\begin{cases} \bar{R}(X) = \{x \mid (\forall x \in M) \wedge ([x]_R \cap X \neq \emptyset)\} = \bigcup \{Y \mid (\forall Y \in M/R) \wedge (Y \cap X \neq \emptyset)\} \\ \text{bn}_R(X) = \bar{R}(X) - \underline{R}(X) \end{cases} \quad (5.21)$$

式中, $[x]_R$ 为对象 x 在等价关系 R 下的等价类, 集合 $\text{bn}_R(X)$ 称为 X 的 R 边界域; $\text{pos}_R(X)$ 称为 X 的 R 正域; $\text{neg}_R(X)$ 称为 X 的 R 负域。设 A 是知识表达系统的属性集, 任意 $\alpha \subseteq A$, 定义: $\text{Mark}(\alpha) = \{0(\alpha \text{ 未被访问}), 1(\alpha \text{ 被访问})\}$ 。

设 $D_T = (M, E \cup G, V, f)$ 是一个决策表, 其中, 论域是对象的一个非空有限集合 $M = \{x_1, x_2, \dots, x_n\}$, $|M| = n$, 则定义 $U_{n \times n}$ 为决策表的差别矩阵:

$$U_{n \times n} = (e_{ij})_{n \times n} = \begin{pmatrix} e_{11} & e_{12} & \cdots & e_{1n} \\ e_{21} & e_{22} & \cdots & e_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ e_{n1} & e_{n2} & \cdots & e_{nn} \end{pmatrix} \quad (5.22)$$

$$e_{ij} = \begin{cases} \{\alpha \mid (\alpha \in \mathbf{E}) \wedge (f_\alpha(x_i) \neq f_\alpha(x_j))\}, & f_G(x_i) \neq f_G(x_j) \\ \emptyset, & f_G(x_i) \neq f_G(x_j) \wedge f_E(x_i) = f_E(x_j) \\ -, & f_G(x_i) = f_G(x_j) \end{cases} \quad (5.23)$$

其中, $i, j = 1, 2, \dots, n$, “-”表示不需要该值或没有该值。

5. 基于数据融合的冶风机故障诊断实例分析

某矿山由于冶风机故障, 污风不能及时排出, 不仅威胁工人的身体健康, 也极大影响正常生产, 使矿山遭受严重的经济损失。经过科学调研后, 该矿领导决定对冶风机运行状态进行远程控制, 在线读取运行参数。但由于影响因素较多, 仍不能从本质上找到导致故障的因素, 故障的表现形式也很不确定, 而且在读取这些参数时发现, 导致相同故障时的运行参数很不一致, 这些正是粗糙集理论研究的对象。运用粗糙集理论从各传感器监测的原始数据出发, 对故障因素进行离散化处理, 确定各条件属性的值, 再依据故障决策属性建立风机故障诊断决策表。该矿山传感器监测的数据经过离散化处理后得到的风机故障诊断决策表, 如表 5-9 所示。 M 表示从监测的数据中随机选出的样本, 即论域; $E = \{E_1, E_2, E_3, E_4\}$ 为条件属性, 分别表示温度、电流、电压、转速 4 个监测参数, 其中, E_1 和 E_2 的值域中 0、1 和 2 分别表示低于正常值、在正常值范围内、高于正常值; E_3 和 E_4 的值域中 0 和 1 分别表示正常值范围之外和正常值范围之内; G 为决策属性, 其中 Y 表示风机正常工作, N 表示风机产生故障。

表 5-9 风机故障诊断决策表

M	E_1	E_2	E_3	E_4	G
1	0	0	0	0	N
2	0	0	0	1	N
3	1	0	0	0	Y
4	2	1	0	0	Y
5	2	2	1	0	Y
6	2	2	1	1	N
7	1	2	1	1	Y
8	0	1	0	0	N
9	0	2	1	0	Y
10	2	1	1	0	Y
11	0	1	1	1	Y
12	1	1	0	1	Y
13	1	0	1	0	Y
14	2	1	0	1	N

由表 5-9 可知, $IND(E) = \{\{1\}, \{2\}, \{3\}, \{4\}, \{5\}, \{6\}, \{7\}, \{8\}, \{9\}, \{10\}, \{11\}, \{12\}, \{13\}, \{14\}\}$; $IND(G) = \{\{1, 2, 6, 8, 14\}, \{3, 4, 5, 7, 9, 10, 11, 12, 13\}\}$ 。显然, $IND(E)$ 属于 $IND(G)$, 所以该风机故障决策表是相容的, 不是矛盾的。为获得风机故障决策表的所有约简, 必须验证以下属性集合: $X_1 = \{E_1, E_4\}$, $X_2 = \{E_1, E_2, E_4\}$, $X_3 = \{E_1, E_3, E_4\}$ 。

(1) 对于属性子集 $X_1 = \{E_1, E_4\}$, 因为 $IND(X_1) = \{\{1, 8, 9\}, \{2, 11\}, \{3, 13\}, \{4, 5, 10\}, \{6, 14\}, \{7, 12\}\}$, 由此得到 $posX_1(G) \{3, 4, 5, 6, 7, 10, 12, 13, 14\} \neq posE(G)$ 。所以, X_1 不是该决策表的 G -约简。

(2) 对于属性子集 $X_2 = \{E_1, E_2, E_4\}$, 因为 $IND(X_2) = \{\{1\}, \{2\}, \{3, 13\}, \{4, 10\}, \{5\}, \{6\}, \{7\}, \{8\}, \{9\}, \{11\}, \{12\}, \{14\}\}$, 由此得到 $posX_2(G) = M = posE(G)$ 。

因此, 风机故障决策表的所有约简为 $REDE(G) = \{\{E_1, E_2, E_4\}, \{E_1, E_3, E_4\}\}$ 。通过属性约简后得到的故障诊断决策表, 如表 5-10 和表 5-11 所示。风机的运转速度和温度是判断风机故障的主要因素, 但不能忽略供电设备电流和电压的变化对风机故障造成的影响。只有融合风机运转速度、温度以及电流或者电压等因素, 才能提高故障诊断的准确率。这个结果与实际情况相符, 验证了运用粗糙集理论和信息融合方法为风机故障进行诊断是一种切实可行的方法。

表 5-10 约简后故障诊断电流属性 X_2 决策表

M	E_1	E_2	E_4	G
1	0	0	0	N
2	0	0	1	N
3	1	0	0	Y
4	2	1	0	Y
5	2	2	0	Y
6	2	2	1	N
7	1	2	1	Y
8	0	1	0	N

续表

M	E_1	E_2	E_4	G
9	0	2	0	Y
10	2	1	0	Y
11	0	1	1	Y
12	1	1	1	Y
13	1	0	0	Y
14	2	1	1	N

表 5-11 约简后故障诊断电压属性 X_3 决策表

M	E_1	E_2	E_4	G
1	0	0	0	N
2	0	0	1	N
3	1	0	0	Y
4	2	0	0	Y
5	2	1	0	Y
6	2	1	1	N
7	1	1	1	Y
8	0	0	0	N
9	0	1	0	Y
10	2	1	0	Y
11	0	1	1	Y
12	1	0	1	Y
13	1	1	0	Y
14	2	0	1	N

5.4.2 基于信息融合的温度场工况分析

1. 信息融合在温度场控制中的应用前景

在智能选冶过程中,特别是在温度场的控制与监测中,信息融合技术显示出巨大的应用潜力。首先,选冶过程是一个复杂的工业流程,涉及多个环节和不同的物理化学反应。这些反应过程通常伴随着热量的释放或吸收,因此准确监测和控制温度场对于保证产品质量、提高生产效率以及确保设备安全至关重要。其次,信息融合技术可以处理不同类型的数据(如数值数据、图像数据、声音数据等),并能够从多个维度对温度场进行分析。

2. 选冶温度场工况分析常见方法

选冶过程通常涉及一系列复杂的物理和化学变化,这些变化往往伴随着能量的交换,因此温度场的控制显得尤为关键。精确的温度场分析对于保证产品质量、提高能效、降低对环境的影响以及确保生产安全至关重要。以下是一些常用的选冶温度场分析方法。

(1) 数值模拟分析。

数值模拟是利用计算机技术和数学模型来模拟实际的选冶过程,它可以根据热传导、热对流和热辐射的基本定律来预测温度场的分布。有限元方法(FEM)和有限差分法(FDM)是两种常用的数值模拟方法。通过建立精确的几何模型和边界条件,可以模拟矿物料在加工过程中的温度变化,从而优化炉子的设计、工艺参数和操作条件。

(2) 实验测量与监测。

实验测量是直接在现场或实验室条件下对温度进行实时监测。使用不同类型的传感器,如热电偶、红外相机和光纤温度传感器等,可以收集温度数据。这些传感器应部署在关键位置,以获取整个选冶系统中的温度分布信息。所收集的数据可以用来验证数值模型的准确性或直接用于过程控制。

(3) 热像仪分析。

热成像技术是通过捕捉物体表面发射的红外辐射来检测温度分布的一种无接触式测量方法。热像仪能够提供实时的二维温度图,有助于快速识别高温区域或热点,这对于预防设备过热和监控冶炼反应非常重要。

3. 实例分析(基于信息融合的高炉料面温度场计算)

1) 理论方法

高炉生产要求边缘煤气流有一定的强度,以减少炉料与炉墙之间的摩擦,防止炉墙结瘤等异常炉况的发生。但是高炉内部环境复杂,目前许多状态难以直接检测,采用的红外图像也不能反映整个料面的温度情况。对高炉内边缘热状态的判断缺乏有效的检测设备和方法。高炉炉墙结构如图 5-15 所示。

高炉料面中部温度场表征了高炉中部区域煤气流的发展状况,是现场操作人员最为关心的内容。通过料面中部温度场的测量可以知道高炉能否顺利运行,运行在什么样的热状态下,是否过热或是过凉,从而能够及时

对炉况进行诊断,有效指导布料操作,保障高炉稳顺高效运行。目前,需获得的高炉料面温度场相关信息主要有以下 3 种。

(1) 十字测温热电偶信息。

十字测温可以检测整个料面上方煤气流的温度。该算法的研究采用热力学原理计算,计算结果的不确定性主要来自炉墙辐射和煤气流的发展方向,而这两方面与料线深度关联很大。当料线深度较小时,对十字测温有辐射的炉墙面积较小,料面各点的煤气流在十字测温位置混合程度较小,因此根据十字测温计算料面温度的可信度较高。当料线深度较大时,炉墙辐射面积加大,料面煤气流在十字测温位置混合程度较大,因此根据十字测温计算料面温度的可信度较低。另外,根据十字测温计算料面温度也与相对位置有关,当料面温度较高的位置煤气流很旺盛时,会直接冲击十字测温梁壁,则测量的可信度较高。相反,料面温度较低,煤气流较弱,可能有一定的混合后再冲击十字测温点,则测量的可信度相对较低。

(2) 炉喉处炉墙中的热电偶信息。

对于大型高炉,其热量传递方式均属于多层平板的稳定导热,这相当于一维平板稳定导热问题的串联叠加。因此,根据炉墙的结构及材料特性,采用热传导计算方法,通过炉墙热电偶检测值,可得到边缘温度估计值。

(3) 边缘矿焦比。

料面对应点的矿焦比反映了高炉轴向矿石和焦炭的比例,该比例是决定整个高炉煤气

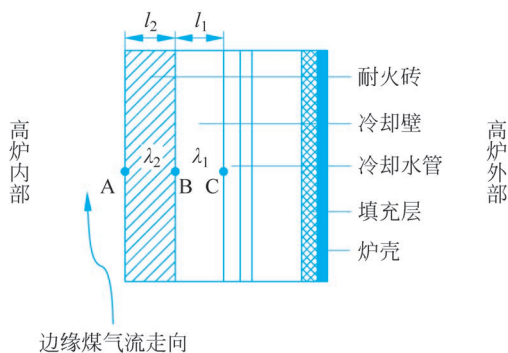


图 5-15 高炉炉墙结构

流分布的最重要的因素,因此它决定了料面温度分布必然趋势和各点温度的相对偏差。这种开环计算方法在计算精度上相对较差,可信度相对较低,但是它能够整体上反映料面不同点的温度发展趋势,具有一定的预测能力。

综上所述,3种料面温度信息各有各的特点,分别从不同的角度计算了料面的温度分布。本实例采用集中式并行卡尔曼滤波方法,将上述3个检测信息进行融合,从而得到料面边缘点的准确温度。在集中式融合结构中,融合中心可以得到所有传感器传送来的原始数据。

2) 仿真实验

首先,将上述3种多源信息依次定义为 $T_E^{CR}, T_E^{WA}, T_E^{OCR}$ 。以料面边缘某点的温度值 T 为 k 时刻的状态方程中的状态 x_k ,以三种检测方法的检测值 T_E^{CR}, T_E^{WA} 和 T_E^{OCR} 为测量方程中 k 时刻的测量值 z_k^1, z_k^2, z_k^3 ,具体如式(5.24)和式(5.25)所示。

$$\begin{cases} x_{k+1} = \mathbf{A}_k x_k + w_k \\ z_{k+1}^i = \mathbf{H}_{k+1}^i x_{k+1} + v_{k+1}^i, i=1,2,3 \end{cases} \quad (5.24)$$

式中, $\mathbf{A}_k$ 是在 $k+1$ 时刻边缘温度的状态转移矩阵; $\mathbf{H}_{k+1}^i$ 为第 i 个传感器在 $k+1$ 时刻的测量矩阵; w_k 和 v_{k+1}^i 表示状态转移过程中的系统噪声和第 i 个传感器的观测噪声,根据温度变化的特点可以假定其是均值为零的白噪声序列,且满足式(5.25)。

$$\begin{cases} \text{cov}(w_k, w_j) = Q_k \delta_{kj}, & Q_k \geq 0 \\ \text{cov}(v_{k+1}^i, v_{j+1}^i) = R_{k+1}^i \delta_{kj}, & R_{k+1}^i > 0 \end{cases} \quad (5.25)$$

其中

$$\delta = \begin{cases} 1, & k=j \\ 0, & k \neq j \end{cases} \quad (5.26)$$

本节采用并行滤波集中式融合算法,首先令

$$\begin{cases} \mathbf{z}_{k+1} = [z_{k+1}^1, z_{k+1}^2, z_{k+1}^3]^T \\ \mathbf{H}_{k+1} = [\mathbf{H}_{k+1}^1, \mathbf{H}_{k+1}^2, \mathbf{H}_{k+1}^3]^T \\ \mathbf{v}_{k+1} = [v_{k+1}^1, v_{k+1}^2, v_{k+1}^3]^T \end{cases} \quad (5.27)$$

根据式(5.28),对应单一传感器量测方程得到融合中心的广义量测方程如下:

$$\mathbf{z}_{k+1} = \mathbf{H}_{k+1} x_{k+1} + \mathbf{v}_{k+1} \quad (5.28)$$

由式(5.25)及式(5.26)的已知条件可得

$$\begin{cases} E(\mathbf{v}_{k+1}) = 0 \\ R_{k+1} = \text{diag}[R_{k+1}^1, R_{k+1}^2, R_{k+1}^3] \end{cases} \quad (5.29)$$

以式(5.24)为融合中心温度状态转移方程,以式(5.28)为融合中心温度广义测量方程。设已知融合中心在 k 时刻的边缘温度状态的融合估计为 $\hat{x}_{k|k}$,相应的误差协方差矩阵为 $\mathbf{P}_{k|k}$,则融合中心相对于3个单一计算方法的集中式融合过程采用卡尔曼滤波器的信息滤波形式,如式(5.30)和式(5.31)所示:

$$\begin{cases} \hat{x}_{k+1|k} = \mathbf{A}_k \hat{x}_{k|k} \\ \mathbf{P}_{k+1|k} = \mathbf{A}_k \mathbf{P}_{k|k} \mathbf{A}_k^T + Q_k \end{cases} \quad (5.30)$$

$$\begin{cases} \hat{x}_{k+1|k+1} = \hat{x}_{k+1|k} + \mathbf{K}_{k+1}(\mathbf{z}_{k+1} - \mathbf{H}_{k+1}\hat{x}_{k+1|k}) \\ \mathbf{K}_{k+1} = \mathbf{P}_{k+1|k+1}\mathbf{H}_{k+1}^T\mathbf{R}_{k+1}^{-1} \\ \mathbf{P}_{k+1|k+1}^{-1} = \mathbf{P}_{k+1|k}^{-1} + \mathbf{H}_{k+1}^T\mathbf{R}_{k+1}^{-1}\mathbf{H}_{k+1} \end{cases} \quad (5.31)$$

式中:

$$\mathbf{R}_{k+1}^{-1} = \text{diag}[(\mathbf{R}_{k+1}^1)^{-1}, (\mathbf{R}_{k+1}^2)^{-1}, (\mathbf{R}_{k+1}^3)^{-1}] \quad (5.32)$$

将式(5.27)和式(5.32)带入式(5.31)中的增益矩阵可得

$$\mathbf{K}_{k+1} = \mathbf{P}_{k+1|k+1}[(\mathbf{H}_{k+1}^1)^T(\mathbf{R}_{k+1}^1)^{-1}, (\mathbf{H}_{k+1}^2)^T(\mathbf{R}_{k+1}^2)^{-1}, (\mathbf{H}_{k+1}^3)^T(\mathbf{R}_{k+1}^3)^{-1}] \quad (5.33)$$

在将式(5.33)和式(5.27)带入式(5.31)中的滤波方程可得

$$\hat{x}_{k+1|k+1} = \hat{x}_{k+1|k} + \mathbf{P}_{k+1|k+1} \sum_{i=1}^3 (\mathbf{H}_{k+1}^i)^T (\mathbf{R}_{k+1}^i)^{-1} (\mathbf{z}_{k+1} - \mathbf{H}_{k+1}^i \hat{x}_{k+1|k}) \quad (5.34)$$

将式(5.27)、式(5.33)和式(5.34)带入式(5.31)中的误差协方差矩阵的逆矩阵可得

$$\mathbf{P}_{k+1,k+1}^{-1} = \mathbf{P}_{k+1|k}^{-1} + \sum_{i=1}^3 (\mathbf{H}_{k+1}^i)^T (\mathbf{R}_{k+1}^i)^{-1} \mathbf{H}_{k+1}^i \quad (5.35)$$

式(5.30)、式(5.34)和式(5.35)构成了基于集中式并行卡尔曼滤波融合料面边缘温度完整的递推方程组。

该递推方程组的初始条件 x 和 $\mathbf{P}$ 由高炉现场操作人员给定,可以证明随着卡尔曼滤波器的工作, x 会逐渐收敛,初始值并不重要。方程组中的状态转移矩阵 $\mathbf{A}$,描述了温度的发展速度,而料面边缘温度的发展速度与边缘的矿焦比 OCR 有关,定义 $\mathbf{A}$ 如式(5.36)所示:

$$\mathbf{A} = 1 + \frac{k}{\text{OCR}} T \quad (5.36)$$

式中, T 为采样时间; OCR 为料面边缘矿焦比; 参数 T 和 k 经过现场测量整定。本书采用的 3 种测量方法直接计算了料面边缘温度,因此测量矩阵 $\mathbf{H}$ 取单位阵 $\mathbf{I}$,过程噪声方差 Q_K 与 OCR 有关,测量方差 $V_K^{\text{CR}}, V_K^{\text{WA}}$ 和 V_K^{OCR} 与料线深度有关。本书根据长时间现场统计确定不同 OCR 和不同料线深度时的过程噪声方程和测量方差。

综上所述,根据集中式并行卡尔曼滤波融合算法可以求解出高炉料面边缘某点的温度值,从而计算出高炉边缘温度场分布。这为判断料面边缘煤气流发展,诊断炉况,以及优化布料操作提供了指导。基于多源信息融合的高炉料面温度场计算方法为选冶生产过程的关键状态检测提供了新的有效手段,通过融合多源信息可以判断煤气流分布,提高高炉炉况诊断的准确性,为选冶过程优化提供坚实的技术基础。