

谢邦昌 朱建平 王小燕 著

厦门大学数据挖掘研究中心
厦门大学管理学院 MBA 中心

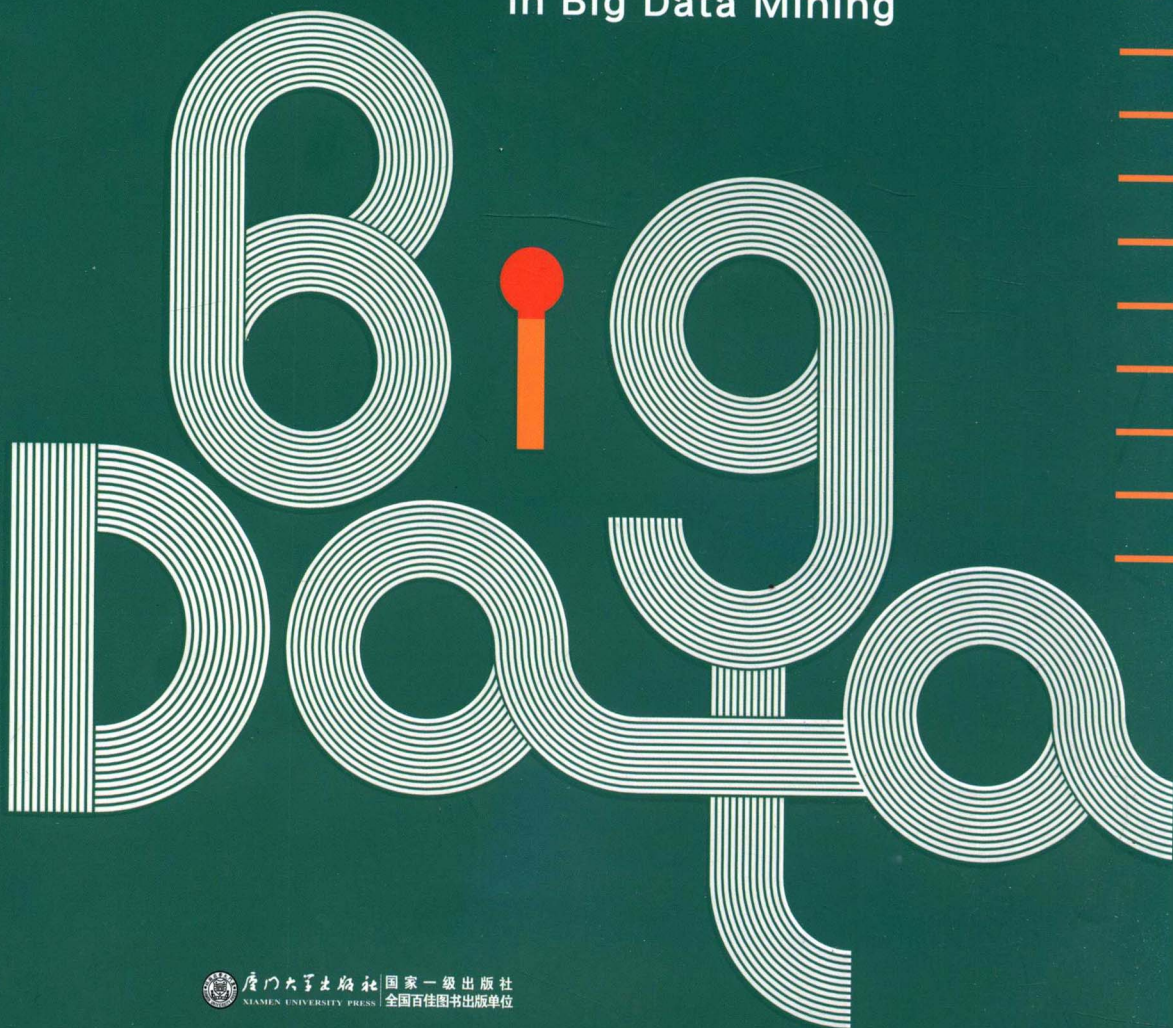
大数据丛书

2

Excel

在大数据挖掘中的应用

The Application of Excel
in Big Data Mining



厦门大学出版社 国家一级出版社
XIAMEN UNIVERSITY PRESS 全国百佳图书出版单位

王小燕 著

Excel

在大数据挖掘中的应用

The Application of Excel
in Big Data Mining



厦门大学出版社
XIAMEN UNIVERSITY PRESS

国家一级出版社
全国百佳图书出版单位

图书在版编目(CIP)数据

Excel 在大数据挖掘中的应用/谢邦昌,朱建平,王小燕著. —厦门:厦门大学出版社, 2016.5
ISBN 978-7-5615-6012-9

I. ①E… II. ①谢…②朱…③王… III. ①数据处理-表处理软件 IV. ①TP274

中国版本图书馆 CIP 数据核字(2016)第 065182 号

出 版 人 蒋东明
责任编辑 吴兴友
封面设计 王 琳
美术编辑 李夏凌
责任印制 吴晓平

出版发行 厦门大学出版社
社 址 厦门市软件园二期望海路 39 号
邮政编码 361008
总 编 办 0592-2182177 0592-2181253(传真)
营销中心 0592-2184458 0592-2181365
网 址 <http://www.xmupress.com>
邮 箱 xmupress@126.com
印 刷 厦门市万美兴印刷设计有限公司

开本 787mm×1092mm 1/16
印张 21.25
插页 2
字数 500 千字
印数 1~3 000 册
版次 2016 年 5 月第 1 版
印次 2016 年 5 月第 1 次印刷
定价 55.00 元

本书如有印装质量问题请直接寄承印厂调换



厦门大学出版社
微信二维码



厦门大学出版社
微博二维码

前言

21 世纪是大数据的时代,又是“互联网+”的时代,借助于计算机存储和分布式计算带来的便捷,信息呈爆炸式增长,各行各业时刻都在生成、收集并处理大量的数据,依赖大数据分析技术的时代悄悄来临。从谷歌推出的流感预测,到百度的市场物价预测和春运的迁徙动态图,能察觉到大数据分析已深入我们的生活。数据挖掘作为大数据分析的核心技术,面对这些繁杂又冗余的大数据,其性能、速度、可视化程度都受到了时代的挑战。

Excel 是当前使用最普遍、入门极快的电子表格软件,它能轻易地完成大量数据的统计分析、处理与制图,不仅功能强大,还操作简单。最新版本 Microsoft Office Excel 2013 是 Excel 发展历程中一个里程碑级的产品,对比 Excel 2010,它非常适用于“大数据”和“互联网+”时代下的数据挖掘工作。它提供了数项让人眼前一新新功能,包括 Power Query, Power View, Power Map 和 PowerPivot。

为了提升 Excel 的数据挖掘和统计分析能力,微软公司提供了一个免费的 SQL Server 2012 数据挖掘外接程序。通过该程序,Excel 2013 能够快速地完成许多专业数据挖掘软件才能完成的任务。在大数据处理非常频繁和迫切的时代,许多学者在学习或工作中面临数据挖掘模型如何实现的难题,因此,我们编写了本书,旨在用最简单的 Excel 2013 软件工具,轻松、快速地实现复杂的数据处理、分析和挖掘工作。

本书的编写以统计原理为主线,以数据挖掘模型的应用为目的。全书共 19 章,基本内容和特点体现为:第 1 篇包括第一章至第四章,概述数据挖掘在大数据时代的重要性,详细总结常用的数据挖掘理论和技术的一般概念、方法技术,并总结了常用的数据挖掘软件和工具的特点和适用性,使读者获得清晰的数据挖掘观

念。第2篇从第五章至第十七章,详细地讲述了数据挖掘加载项的配置要求和安装过程。接着从理论、操作、实例展示三个方面重点讲述了数据挖掘加载项的9种数据挖掘模型的使用,包括决策树、贝叶斯概率分类、关联规则、聚类分析、时序聚类、线性回归、Logistic回归、神经网络聚类、时间序列分析。内容全面,涵盖了目前学术界和实业界主要的数据挖掘技术和方法。同时,通过实例操作介绍多种数据挖掘辅助技术,包括关键影响因素、类别检测、从示例填充、预测、异常值检测、应用场景分析。第3篇为实例篇,包含第十八、十九章,以台湾房屋信用贷款数据和台湾健康食品行业数据为例,应用第2篇中多种数据挖掘方法,在Excel 2013数据加载项上展开实例分析,通过详细的操作演示和结果解释,读者可获得实际的数据挖掘经验,并能迅速在自己的学习和工作中加以应用。最后为附录部分,通过详细的操作示范,讲述了Excel 2013的四大特色数据分析功能,包括Power Query,Power View,Power Map和PowerPivot。

本书适合多层次多专业的读者,如数据、统计、经济金融、工商管理等专业的大学生、研究生学习,还适合于非统计类从事相关数据分析的人员阅读。

本书由台北医学大学谢邦昌教授设计整体框架,并同厦门大学朱建平教授、湖南大学王小燕助理教授共同撰写。在本书编写过程中,厦门大学研究生范新妍、黄晓雯、张悦涵为本书第2篇部分章节和附录的资料收集和录入,以及部分案例的整理做了大量的工作。谢邦昌和朱建平教授对本书进行了修改和总纂,台北医学大学邓光甫教授审阅了全书。

本书在编写和出版过程中,得到了厦门大学数据挖掘研究中心、台北医学大学大数据研究中心、厦门大学管理学院MBA中心、湖南大学金融与统计学院和厦门大学出版社的支持,陈丽贞和吴兴友同志为本书的组稿、编辑做了大量的工作,在此表示衷心感谢!编写一本好的书并不容易,尽管我们努力想奉献给读者一本满意的书,但仍有达不到读者各方面的要求的地方。书中难免有疏漏或错误之处,恳请读者多提宝贵意见,以便今后进一步修改与完善。

本书的编写和出版得到了国家社会科学基金重大项目“大数据与统计学理论的发展研究”(13&2D148)的资助。

作者

2016年1月

BIG
DATA

第一篇 基本知识

003 第一章 大数据与数据挖掘

- 1.1 大数据的定义 / 003
- 1.2 大数据的 4V 特征 / 003
- 1.3 大数据的预测魅力 / 004
- 1.4 数据挖掘定义 / 005
- 1.5 数据挖掘的重要性 / 006
- 1.6 数据挖掘功能 / 006
- 1.7 数据挖掘步骤 / 007
- 1.8 数据挖掘建模的标准 CRISP-DM / 007
- 1.9 大数据时代数据挖掘面临的挑战 / 009

010 第二章 数据挖掘运用理论及技术

- 2.1 回归分析 / 010
 - 2.1.1 简单线性回归分析 / 010
 - 2.1.2 多元回归分析 / 010
 - 2.1.3 岭回归分析 / 011
 - 2.1.4 Logistic 回归分析 / 012
- 2.2 关联规则 / 012
- 2.3 聚类分析 / 012
- 2.4 判别分析 / 014
- 2.5 神经网络分析 / 015
- 2.6 决策树分析 / 017
- 2.7 其他分析方法 / 018

020 第三章 数据挖掘与其他相关领域的关系

- 3.1 数据挖掘与统计分析的不同 / 020
- 3.2 数据挖掘与数据仓储的关系 / 020
- 3.3 KDD 与数据挖掘的关系 / 022
- 3.4 OLAP 与数据挖掘的关系 / 022
- 3.5 数据挖掘与机器学习的关系 / 023
- 3.6 网络信息挖掘和数据挖掘有什么不同? / 023

025 第四章 数据挖掘商业软件产品及其应用现状

- 4.1 数据挖掘工具分类 / 025
- 4.2 各工具的简介 / 025
- 4.3 客户关系管理(CRM) / 026
- 4.4 数据挖掘在各行业的应用 / 027

第二篇 Excel 2013 数据挖掘模型**031 第五章 安装与配置 Excel 2013 数据挖掘加载项**

- 5.1 系统需求 / 031
- 5.2 开始安装 / 031
- 5.3 完成安装验证 / 035
- 5.4 配置设定 / 035
- 5.5 设定完成验证 / 040

042 第六章 Excel 2013 数据挖掘入门

- 6.1 Excel 2013 数据挖掘工具栏介绍 / 042
- 6.2 数据挖掘使用说明 / 042
 - 6.2.1 目录 / 043
 - 6.2.2 使用者入门 / 043
 - 6.2.3 视频和教程 / 044
- 6.3 数据挖掘连接设定 / 044
 - 6.3.1 设定目前的连接 / 044
 - 6.3.2 跟踪 / 046
- 6.4 数据准备 / 47
 - 6.4.1 浏览数据 / 047
 - 6.4.2 清除数据 / 050
 - 6.4.3 分割数据 / 053
- 6.5 数据建模 / 057
- 6.6 准确性和验证 / 058
 - 6.6.1 准确性图表 / 058
 - 6.6.2 分类矩阵 / 059
 - 6.6.3 利润图 / 059
- 6.7 模型用法 / 060

- 6.7.1 浏览 / 061
- 6.7.2 查询 / 063
- 6.8 模型管理 / 064
 - 6.8.1 重命名挖掘结构 / 064
 - 6.8.2 删除挖掘结构 / 065
 - 6.8.3 清除挖掘结构 / 065
 - 6.8.4 使用原始数据处理挖掘结构 / 065
 - 6.8.5 用新数据处理挖掘结构 / 065
 - 6.8.6 导出挖掘结构 / 067
 - 6.8.7 导入挖掘结构 / 067

068 第七章 决策树

- 7.1 基本概念 / 068
- 7.2 决策树模块的建立:三种形式 / 068
- 7.3 决策树与判别函数比较 / 068
- 7.4 计算方法 / 069
 - 7.4.1 制定预测精确性的标准规范 / 069
 - 7.4.2 选择分裂(分层)技术 / 070
 - 7.4.3 定义停止分裂(分层)的时间点 / 070
 - 7.4.4 选择适当大小的决策树 / 071
- 7.5 Excel 2013 决策树算法操作步骤 / 071

082 第八章 贝叶斯概率分类

- 8.1 基本概念 / 082
- 8.2 Excel 2013 贝叶斯概率分类操作步骤 / 085

095 第九章 关联规则

- 9.1 基本概念 / 095
- 9.2 关联规则的种类 / 096
- 9.3 关联规则的算法:Apriori 算法 / 097
 - 9.3.1 执行步骤 / 097
 - 9.3.2 优点 / 097
 - 9.3.3 缺点 / 097
- 9.4 Excel 2013 关联规则操作步骤 / 097

103 第十章 聚类分析

- 10.1 基本概念 / 103
 - 10.1.1 搜集数据 / 103
 - 10.1.2 转换成相似矩阵 / 103
- 10.2 层次聚类分析(hierarchical clustering methods) / 104
 - 10.2.1 系统聚类法 / 104
 - 10.2.2 逐步聚类法 / 104
 - 10.2.3 逐步分解法 / 104
 - 10.2.4 有序样本的聚类 / 104
- 10.3 聚类分析原理 / 105
- 10.4 Excel 2013 聚类分析操作步骤 / 109

126 第十一章 时序聚类

- 11.1 基本概念 / 126
- 11.2 相关研究 / 126
- 11.3 Excel 2013 时序聚类操作步骤 / 128

138 第十二章 线性回归

- 12.1 基本概念 / 138
- 12.2 简单回归分析 / 139
- 12.3 多元回归分析 / 142
- 12.4 回归变量的选择 / 144
- 12.5 Excel 2013 线性回归操作步骤 / 145

155 第十三章 Logistic 回归

- 13.1 基本概念 / 155
- 13.2 logit 变换 / 156
- 13.3 Logistic 分布 / 157
- 13.4 2×2 表的 Logistic 回归模型 / 158
- 13.5 Excel 2013 Logistic 回归操作步骤 / 159

174 第十四章 神经网络

- 14.1 基本概念 / 174
- 14.2 神经网络的特性 / 176
- 14.3 神经网络的架构与训练算法 / 176

- 14.4 神经网络应用 / 177
- 14.5 神经网络的优缺点 / 178
- 14.6 神经网络的限制 / 178
- 14.6 Excel 2013 神经网络操作步骤 / 179

196 第十五章 时间序列分析

- 15.1 基本概念 / 196
- 15.2 时间序列的成分 / 199
 - 15.2.1 趋势成分 / 199
 - 15.2.2 循环成分 / 199
 - 15.2.3 季节成分 / 200
 - 15.2.4 不规则成分 / 200
- 15.3 利用修匀法预测 / 200
 - 15.3.1 移动平均 / 201
 - 15.3.2 加权移动平均 / 202
 - 15.3.3 指数修匀 / 203
- 15.4 用趋势投射预测时间序列 / 205
- 15.5 预测含趋势与季节成分的时间序列 / 206
 - 15.5.1 消除时间序列的季节性 / 206
 - 15.5.2 消除季节性的时间序列, 辨识趋势 / 206
 - 15.5.3 循环成分 / 207
- 15.6 利用回归模型预测时间序列 / 207
- 15.7 其他预测模型 / 208
- 15.8 单变量时间序列预测模型 / 209
 - 15.8.1 自回归模型 (autoregressive models, AR model) / 210
 - 15.8.2 移动平均过程模型 (moving average process model, MA model) / 210
 - 15.8.3 AR-MA 模型 (mixed AR-MA model) / 210
 - 15.8.4 季节循环性时间序列模型 / 210
- 15.9 时间趋势预测模型 / 211
- 15.10 Excel 2013 时间序列操作步骤 / 212

217 第十六章 DMX 介绍

- 16.1 DMX 介绍 / 217
- 16.2 DMX 函数介绍 / 219
 - 16.2.1 模型建立 / 219
 - 16.2.2 模型训练 / 220



- 16.2.3 模型使用(预测) / 220
- 16.2.4 其他函数语法 / 221
- 16.3 DMX 数据挖掘语法 / 224
 - 16.3.1 决策树 / 225
 - 16.3.2 贝叶斯概率分类 / 226
 - 16.3.3 关联规则 / 226
 - 16.3.4 聚类分析 / 227
 - 16.3.5 时序聚类 / 228
 - 16.3.6 线性回归分析 / 229
 - 16.3.7 Logistic 回归 / 230
 - 16.3.8 神经网络 / 231
 - 16.3.9 时间序列 / 232
- 16.4 DMX 应用范例 / 233
 - 16.4.1 分类(classification) / 233
 - 16.4.2 估计(estimation) / 234
 - 16.4.3 预测(prediction) / 235
 - 16.4.4 关联分组(affinity grouping) / 236
 - 16.4.5 同质分组(clustering) / 237

238 第十七章 其他分析工具

- 17.1 分析关键影响因素 / 238
- 17.2 检测类别 / 242
- 17.3 从示例填充 / 245
- 17.4 预测 / 246
- 17.5 突出显示异常值 / 248
- 17.6 应用场景分析 / 250
 - 17.6.1 目标查找 / 251
 - 17.6.2 假设 / 253

第三篇 实例

257 第十八章 台湾房屋信用贷款违约状况分析

- 18.1 数据说明 / 257
- 18.2 建立模型 / 258

267 第十九章 台湾健康食品行业分析

19.1 数据说明 / 267

19.2 建立模型 / 269

19.2.1 准确性图表(预测列“违约注记”=1) / 269

19.2.2 分类矩阵 / 271

附 录**277 附录一 Power Query 简介**

1.Power Query 简介 / 277

2.Power Query 安装 / 277

3.Power Query 进行数据分析实例 / 279

292 附录二 Power View 简介

1.Power View 简介 / 292

2.Power View 操作 / 292

301 附录三 Power Map

1.Power Map 简介 / 301

2.数据要求 / 301

3.安装 Power Map / 302

4.Power Map 的功能实现 / 303

310 附录四 PowerPivot

1.PowerPivot 加载 / 310

2.向 PowerPivot 工作簿中加载数据 / 312

3.在数据之间创建关系 / 318

4.创建计算列 / 321

5.创建数据透视表 / 325

第一篇

BIG
DATA

基本知识

BIG DATA



第一章 大数据与数据挖掘

随着云计算、电子商务、网络技术的飞速发展,数据种类和规模在金融服务、零售、科研、医疗保健、能源、交通等诸多行业飞速增长。2006 年全球新产生的数据量大约为 180 EB,2011 年达到 1.8 ZB,有研究机构预测到 2020 年这个数字将增长到 35.2 ZB,大数据时代伴随而来,它成为继“云端”以后在 IT 界深入人心的流行词汇。

1.1 大数据的定义

“大数据”一词最早出现在 1997 年迈克尔·考克斯的论文中,直到 2008 年 9 月,《科学》杂志发表论文“Big Data: Science in the Petabyte Era”,大数据一词开始广泛传播。如同“云端”一样,在出现之初并没有给它明确的定义。大数据到底是什么,它极其复杂,至今还没有一个得到普遍认同的官方定义。

大数据不仅仅等于数据量大,亚马逊大数据科学家 John Rauser 将其定义为“任何超过一台计算机处理能力的庞大数据量”;Informatica 中国区首席产品顾问但彬认为“大数据=海量数据+复杂类型的数据”,其规模和复杂程度超过了常用技术按照合理的成本和时限捕捉、管理以及处理这些数据集的能力;维基百科将其描述为“任何由于规模庞大且高度复杂而难以通过现有数据库管理工具或者传统数据处理应用进行处理的数据集”。

通俗地说,大数据是现有的一般技术难以管理的海量数据。“现有的一般技术难以管理”,一方面是由于大数据的结构极其复杂,传统的关系数据无法管理;另一方面是大数据在量方面的急剧增加,导致查询数据的反应时间超过容许范围。

1.2 大数据的 4V 特征

顾名思义,大数据具有数据量大的特征,但这不是它的全部,一般而言,它具有以下四个特征:

大量性(volume):是指数据量的庞大性以及规模的完整性,主要体现在存储和计算两方面,其单位从TB上升到PB级别,在互联网时代,社交网络、电子商务等把人类带入了“PB”新时代。这与数据存储和网络技术的发展密切相关。据统计,互联网一天产生的全部内容足以制作1.68亿张DVD,一天发出2940亿封邮件以及200万个帖子。

多样性(variety):它不仅仅具有量的剧增性,同时也有数据复杂性的提升,主要是指大数据包含的数据类型繁多。电商和物联网的发展,使得它不再局限于传统的结构化数据,还包括半结构以及非结构化数据,如文本、音频、图片、搜索记录等,后者所占比例高达85%。

高速性(velocity):主要体现在大数据的增长速度快,以及传输和处理速度快。大数据通过云计算,可以在短短几十分钟内,将传统模式下十几天才能完成的数据存储完毕。

价值性(value):这主要体现了大数据的应用价值。但是它的价值性具有稀疏性、不确定性和多样性。数据量大并不意味着它们都是有效的,这其中大部分数据都是噪音,价值密度很低。数据量增长的同时,隐藏的有价值信息并没有相应比例增长,从而获取有用信息的难度加大。

1.3 大数据的预测魅力

大数据的预测目前主要体现在商业价值上,已有相关的公司利用大数据技术进行研究。国内最有代表性的有三家公司,分别是百度、阿里巴巴以及腾讯。具体来说,大数据应用到了如下方面:

(1) 大数据预测世界杯赛事

在世界杯紧张进行的同时,百度、谷歌、微软以及高盛等科技公司都推出了赛事预测平台,百度预测结果尤为亮眼。在小组赛阶段,百度以58.33%的准确率居首,在淘汰赛阶段,百度和微软全程预测16进8及8进4的比赛,准确率100%;谷歌预测了12场比赛仅一场错误,准确率为91.67%。百度在这场预测大赛中领跑,其大数据研究院利用百度大数据搜索过去5年内全球987支球队的3.7万场比赛数据,并与乐彩网等建立数据战略合作伙伴关系,将博彩数据融入预测模型,同时考虑了团队实力、主场优势、最近表现、世界杯整体表现和博彩公司赔率等五个因素。这意味着大数据已取代章鱼保罗来掌控赛事预测。

(2) 百度迁徙

百度作为技术型公司,已诞生了多个大数据背景下的产品。百度迁徙是百度在2014年春运期间基于大数据技术推出的一个品牌项目。它基于LBS大数据进行全样本的数据处理、分析和挖掘,在业界首次实现了全程、动态、即时和直观地展示人口迁徙的轨迹与特征,并实现可视化。

(3) 疾病预测

基于用户的搜索数据以及购物行为,并结合天气、环境和人口流动等因素,百度与中

国疾病预防控制中心合作,就流感、肺炎、肺结核、性病这四种疾病,对全国每一个省份以及大多数市区的活跃度和趋势图等进行了全面的监控。这是继世界杯预测后百度大数据预测的又一款产品。用户可以查看过去 30 天以内的数据以及未来 7 天的预测趋势。

(4) 市场物价预测

CPI 表征物价的浮动情况,但是目前的 CPI 存在偏差性,仅仅编制一个物价总指数显然不足以准确反映不同收入阶层的实际价格水平。同时,CPI 存在较长的时滞性,会弱化它的信息反馈功能。大数据时代数据的获取不再是问题,数据更新速度更快,这为统计调查提供了海量的原始资料。大数据可以帮助人们了解未来物价走向,提前预知通货膨胀或经济危机。阿里巴巴团队通过阿里 B2B 大数据提前预知亚洲金融危机,这是市场物价预测的典型案例。

此外,阿里巴巴在业务驱动下,已通过大数据着手分析各种用户消费行为,“云端+大数据”已成为它的重要战略。阿里巴巴会记录淘宝和天猫的每一笔消费记录,交易及信用数据成为它的一手资料,从而建立其数据库。而云平台阿里云是阿里巴巴支撑大数据的主要组成部分,在网购高峰期如 2013 年的“双十一”,75% 的交易都是在阿里云平台上运行。

大数据的预测应用还有景点预测、城市旅游预测、高考考研预测、金融预测、票房预测、选举预测、奥斯卡获奖预测等等,大数据的预测魅力将在各个领域得到体现。

1.4 数据挖掘定义

“Data mining is the process of seeking interesting or valuable information in large data bases.”

数据挖掘(data mining)是近年来数据库应用领域中相当热门的话题。数据挖掘指在数据库中,利用各种分析方法对累积的大量历史数据进行分析、归纳与整合等工作,以提取出有用的信息,找出有意义且用户有兴趣的模式(interesting patterns),供企业管理层进行决策。

数据挖掘指找寻隐藏在数据中的信息,如趋势(trend)、模式(pattern)及相关性(relationship)的过程,也就是从数据中发掘信息或知识(也称 knowledge discovery in databases, KDD),也有人称为数据考古学(data archaeology)、数据模式分析(data pattern analysis)或功能相依分析(functional dependency analysis),目前已被许多研究人员视为结合数据库系统与机器学习技术的重要领域,许多产业界人士也认为此领域是一项增加各企业潜能的重要指标。

事实上,数据挖掘并不只是一种技术或一套软件,而是一种结合多项专业知识的应用。但我们对数据挖掘应有一个正确的认识:数据挖掘工具是从数据中发掘出各种可能的假设(hypothesis),但并不帮你证实这些假设,也不帮你判断这些假设对你的价值。