

资深数据分析咨询师多年经验结晶，内容全面而深入，为高效利用各类工具和算法进行数据分析提供最佳指导

通过大量典型数据分析案例，全面阐释分类分析、聚类分析、数据可视化及预测方面的各种技术和方法，为快速掌握并灵活运用数据分析技术提供最佳实践指南



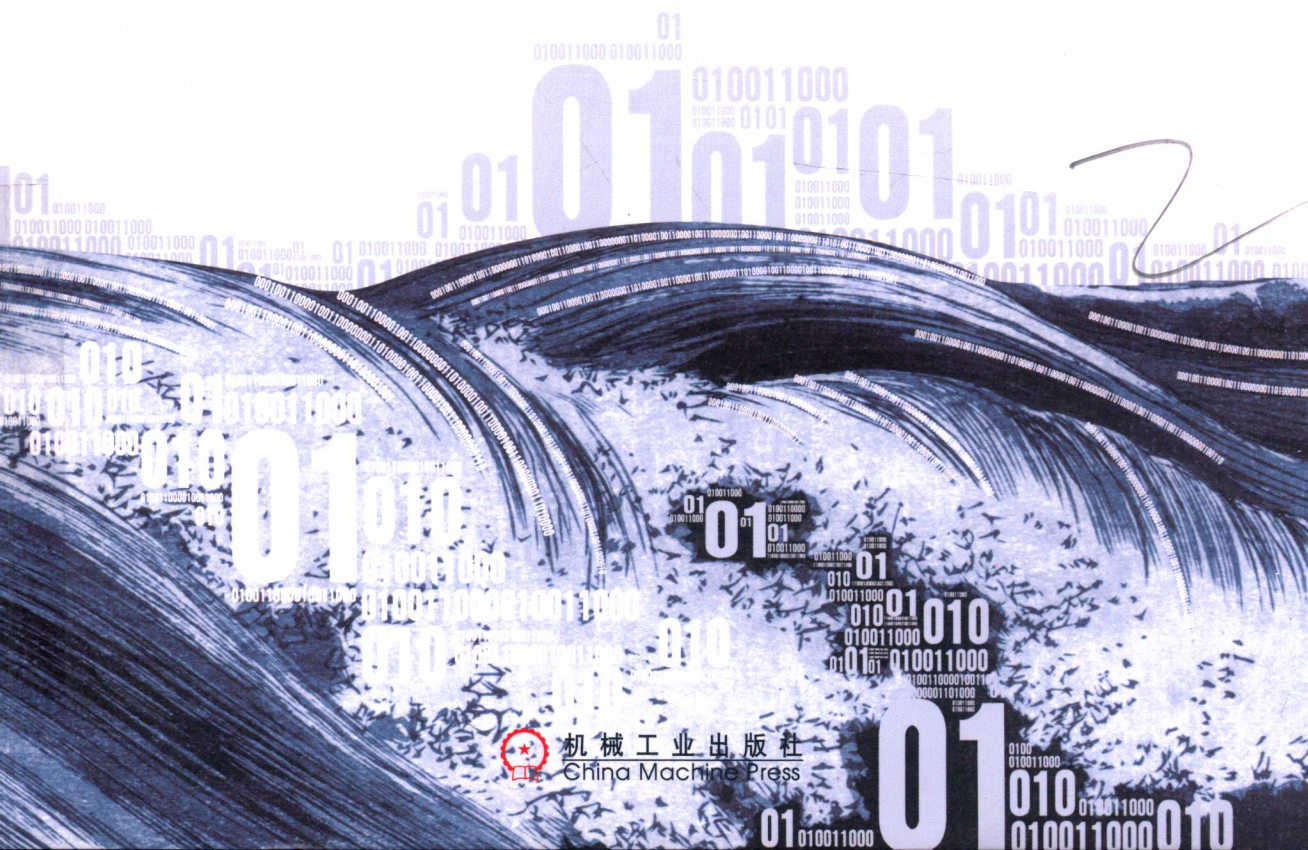
技术丛书

Practical Data Analysis, Second Edition

实用数据分析

(原书第2版)

[美] 赫克托·奎斯塔 (Hector Cuesta) 著
桑帕斯·库马尔 (Dr. Sampath Kumar)
刁晓纯 译



机械工业出版社
China Machine Press



技术丛书

Practical Data Analysis, Second Edition

实用数据分析

(原书第2版)

[美] 赫克托·奎斯塔 (Hector Cuesta) 著
桑帕斯·库马尔 (Dr. Sampath Kumar)

刁晓纯 译



机械工业出版社
China Machine Press

图书在版编目 (CIP) 数据

实用数据分析 (原书第 2 版) / (美) 赫克托·奎斯塔 (Hector Cuesta), (美) 桑帕斯·库马尔 (Dr. Sampath Kumar) 著; 刁晓纯译. —北京: 机械工业出版社, 2017.8
(大数据技术丛书)

书名原文: Practical Data Analysis, Second Edition

ISBN 978-7-111-57921-2

I. 实… II. ①赫… ②桑… ③刁… III. 统计数据—统计分析 IV. O212.1

中国版本图书馆 CIP 数据核字 (2017) 第 216292 号

本书版权登记号: 图字: 01-2017-0487

Hector Cuesta, Dr. Sampath Kumar: *Practical Data Analysis, Second Edition* (ISBN: 978-1-78528-971-2).
Copyright © 2016 Packt Publishing. First published in the English language under the title “Practical Data Analysis, Second Edition”.

All rights reserved.

Chinese simplified language edition published by China Machine Press.

Copyright © 2017 by China Machine Press.

本书中文简体字版由 Packt Publishing 授权机械工业出版社独家出版。未经出版者书面许可, 不得以任何方式复制或抄袭本书内容。

实用数据分析 (原书第 2 版)

出版发行: 机械工业出版社 (北京市西城区百万庄大街 22 号 邮政编码: 100037)

责任编辑: 陈佳媛

责任校对: 李秋荣

印 刷: 北京诚信伟业印刷有限公司

版 次: 2017 年 9 月第 1 版第 1 次印刷

开 本: 186mm×240mm 1/16

印 张: 15.5

书 号: ISBN 978-7-111-57921-2

定 价: 59.00 元

凡购本书, 如有缺页、倒页、脱页, 由本社发行部调换

客服热线: (010) 88379426 88361066

投稿热线: (010) 88379604

购书热线: (010) 68326294 88379649 68995259

读者信箱: hzit@hzbook.com

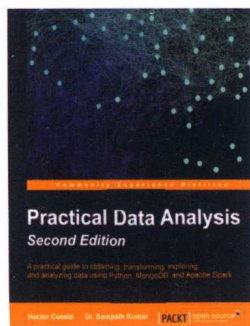
版权所有·侵权必究

封底无防伪标均为盗版

本书法律顾问: 北京大成律师事务所 韩光 / 邹晓东

内容简介

本书共15章：第1章探讨数据分析的基本原理和数据分析步骤；第2章解释如何清洗并准备数据；第3章展示在JavaScript可视化框架下应用D3.js来实现各类数据的可视化方法；第4章介绍如何应用朴素贝叶斯算法来区分垃圾邮件；第5章讲解应用动态时间规整方法寻找图像间的相似性；第6章介绍使用随机游走算法和可视化的D3.js动画技术模拟股票价格；第7章介绍核岭回归（KRR）的原理以及应用；第8章描述如何使用支持向量机方法进行分类分析；第9章介绍应用细胞自动机方法对传染病进行建模；第10章解释如何应用Gephi从Facebook获取社交媒体图谱并实现可视化；第11章介绍如何应用Twitter数据进行情感分析；第12章介绍如何使用MongoDB进行数据处理和聚合；第13章详细介绍如何在MongoDB数据库中应用MapReduce编程模型；第14章介绍如何应用Jupyter和Wakari开展线上数据分析；第15章介绍如何使用Apache Spark处理数据。



原书封面

作者简介

赫克托·奎斯塔(Hector Cuesta)

Dataxios（一家机器智能研发公司）的创始人及首席数据科学家，拥有信息学士及计算机科学硕士学位。他在金融、零售、金融科技、在线学习、人力资源等领域提供数据驱动产品设计的咨询服务。可以关注他的推特：
<https://twitter.com/hmCuesta>。

桑帕斯·库马尔(Dr. Sampath Kumar)

Telangana 大学应用统计系的助理教授和系主任，拥有5年研究生教学经验，有超过4年的工作经验。他是SAS和MATLAB软件高级编程人员，擅长利用SPSS、SAS、R、Minitab、MATLAB等软件进行数据统计。他在不同的应用学科和纯统计专业，如预测建模、应用回归分析、多变量数据分析、运营管理等方面具有教学经验。

The Translator's Words 译者序

早在 2013 年 7 月，因为参加了“数据驱动企业分析变革商业”2013 第三届大数据世界论坛，认识了机械工业出版社的编辑王春华。当时从事数据方面的工作已经 3 年多，身处国内大型央企，我所涉及的数据分析工作非常广泛，既跨越了银行保险、公共服务、电子商务及速递物流等行业，也包含了对客户、渠道、价格、实物网效率、经营业绩等多方面的分析。在当时的工作中遇到了很多问题，既有组织方面的，也有方法效率方面的，王编辑推荐我阅读本书第 1 版原著，并且我也有幸参与翻译了部分章节。令人惊喜的是，书中所介绍的广泛案例、先进的方法以及诸多便利的工具都对数据分析工作有很多帮助和值得借鉴的地方。

针对本书，我的主要体会会有三方面：

第一，本书包含丰富的案例。书中介绍的案例涉及垃圾邮件的分类分析、图像匹配、流行病暴发事件分析、社交网络的数据获取和分析、对文本型数据进行情感分析、股票价格以及黄金价格走势分析等。

第二，本书所涉内容包含了数据分析的全流程，包括了数据准备和处理、多类型建模、数据可视化展示等。初次接触数据分析的读者可以由浅入深地了解分析的全貌。

第三，本书充分体现了大数据的特点，既介绍了对结构化数据的处理也介绍了对非结构化数据的处理，数据类型丰富。书中所涉数据包括时间序列数据、数值型数据、多维度数据和社交媒体数据、文本型数据等多种形式，可以帮助读者获得对数据分析的真知灼见。

时隔几年，机械工业出版社联系上我，询问我是否愿意翻译本书第 2 版，我二话不说接下了这个任务，这几年随着数据工作方面的积累，对于本书，除了有更深的体会，也重新回顾、整理了当年翻译的内容。随着“大数据”技术的发展，本书最后一章也新增了对 Cloudera VM 和 Apache Spark 的介绍，使读者了解其在大数据领域的地位，并掌握一些常见的方法和操作。这又是一次温故而知新的历程。

书中部分内容是按照原文直译的，难免有不完整或者偏颇的地方，请读者批评指正，也欢迎广大读者与我交流沟通，我的邮箱是 jacqueline_dut@hotmail.com。

刁晓纯

2017 年 6 月

作者简介 *About the Author*

Hector Cuesta Dataxios (一家机器智能研发公司) 的创办人及首席数据科学家, 拥有信息学士及计算机科学硕士学位。他在金融、零售、金融科技、在线学习、人力资源等领域提供数据驱动产品设计的咨询服务。在空闲时间, 他热衷于研究机器人。可以关注他的推特: <https://twitter.com/hmCuesta>。

本书献给我的妻子 Yolanda 和我可爱的孩子 Damian 和 Issac, 他们为我的生活带来了无比的快乐。同时把本书献给我的父母 Elena 和 Miguel, 感谢他们对我的支持和爱护。

Dr. Sampath Kumar Telangana 大学应用统计系的助理教授和系主任, 他拥有理学硕士、哲学硕士和统计学博士学位, 拥有 5 年研究生教学经验, 有超过 4 年的工作经验。他是 SAS 和 MATLAB 软件高级程序员, 专长是利用 SPSS、SAS、R、Minitab、MATLAB 等软件进行数据统计。他在不同的应用学科和纯统计专业 (如预测建模、应用回归分析、多变量数据分析、运营管理等) 方面具有教学经验。

About the Reviewers 审校者简介

Chandana N. Athauda 目前是 BAG (Brunei Accenture Group) 的员工，他在 Brunei 是一名技术顾问。他主要关注商务智能、大数据和数据可视化工具技术。他已经在 IT 行业从业超过 15 年（曾获前微软最有价值员工和微软 TFS 管理员）。他对 IT 行业充满了工作热情，他的工作职业从程序员贯穿到技术顾问。

如果有兴趣与 Chandana 讨论本书，请发送邮件到 info@inzeek.net 或是上推特 @inzeek。

Mark Kerzner 大数据架构师及培训师。他是 Elephant Scale 的创始人及负责人，这家企业为不同领域提供大数据的顾问咨询及培训。同时他也是《HBase Design Patterns》一书的作者。

我要感谢我的联合创始人 Sujee Maniyam 和他的同事 Tim Fox，还要感谢所有的老师及学生。最后同样重要的是，感谢家人对我的帮助。

前言 *Preface*

本书提供了一系列将数据转化为重要结论的现实案例。书中覆盖了广泛的数据分析工具和算法，用于进行分类分析、聚类分析、数据可视化、数据模拟以及预测。本书旨在帮助读者了解数据从而找到相应的模式、趋势、相互关系以及重要结论。

书中所包括的实用项目充分利用了 MongoDB、D3.js 和 Python 语言，并采用代码片段和详细描述的方式呈现本书的核心概念。

本书主要内容

第 1 章探讨数据分析的基本原理和数据分析步骤。

第 2 章解释如何清洗并准备好数据来开展分析，同时介绍数据清洗工具 OpenRefine 的使用方法。

第 3 章展示在 JavaScript 可视化框架下应用 D3.js 语言来实现各类数据的可视化方法。

第 4 章介绍应用朴素贝叶斯 (Naive Bayes) 算法来区分垃圾文本的一种二元分类法。

第 5 章展示一个应用动态时间规整方法来寻找图像间相似性的项目。

第 6 章解释如何使用随机漫步算法和可视化的 D3.js 动画技术来模拟股票价格。

第 7 章介绍核岭回归 (Kernel Ridge Regression, KRR) 的原理以及如何使用此方法和时间序列数据来预测黄金价格。

第 8 章描述如何使用支持向量机的方法进行分类分析。

第 9 章介绍对流行病进行模拟计算的基本概念并解释如何应用细胞自动机方法、D3.js 和 JavaScript 语言来模拟流行病爆发。

第 10 章解释如何应用 Gephi 从 Facebook 获取社交媒体图谱并使之实现可视化。

第 11 章解释如何应用 Twitter 的应用程序编程接口 (API) 来获取 Twitter 的数据。读者也将看到如何改进文本分类分析方法并将其应用于情感分析。这一过程在自然语言工具包 (Natural Language Toolkit, NLTK) 中应用了朴素贝叶斯算法。

第 12 章介绍在 MongoDB 数据库中进行基本操作以及分组、过滤和聚合的方法。

第 13 章详细介绍如何在 MongoDB 数据库中应用 MapReduce 编程模型。

第 14 章解释如何使用 Wakari 平台，同时介绍在 IPython 中运用 pandas 进行数据处理和使用 PIL 图像处理库的方法。

第 15 章介绍如何在 Cloudera VM 上使用分布式文件系统及数据环境。最后，利用实际案例介绍 Apache Spark 的主要特征。

阅读准备

使用本书需要掌握如下技术：

- ☐ Python
- ☐ OpenRefine
- ☐ D3.js
- ☐ mlpy
- ☐ NLTK
- ☐ Gephi
- ☐ MongoDB

读者对象

本书主要面向那些希望能够实际开展数据分析和数据可视化的软件开发人员、分析人员、计算机科学家。同时，本书也希望能够为读者提供包含时间序列数据、数值型数据、多维度数据和社交媒体数据、文本型数据等多种数据形式的实际案例，以帮助读者获得对数据分析的真知灼见。

读者不需要具备数据分析的经验，但仍需要对统计学和 Python 编程有基础性的了解。

下载本书相关资源

读者可登录华章网站 (<http://www.hzbook.com>) 下载本书的相关资源。

目 录 *Contents*

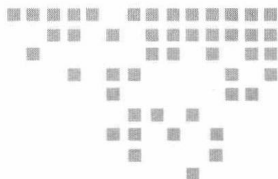
译者序	
作者简介	
审校者简介	
前言	
第1章 开始	1
1.1 计算机科学	1
1.2 人工智能	2
1.3 机器学习	2
1.4 统计学	2
1.5 数学	2
1.6 专业领域知识	3
1.7 数据、信息和知识	3
1.7.1 数据、信息和知识之间的相互性	3
1.7.2 数据的本质	4
1.8 数据分析过程	5
1.8.1 问题	6
1.8.2 数据准备	6
1.8.3 数据探索	7
1.8.4 预测建模	7
1.8.5 结果可视化	8
1.9 定量与定性数据分析	9
1.10 数据可视化的重要性	9
1.11 大数据	10
1.12 自我量化	12
1.12.1 传感器和摄像头	12
1.12.2 社交网络分析	13
1.13 本书的工具和练习	13
1.13.1 为什么使用 Python	14
1.13.2 为什么使用 mlpy	14
1.13.3 为什么使用 D3.js	14
1.13.4 为什么使用 MongoDB	15
1.14 小结	15
第2章 数据预处理	16
2.1 数据源	16
2.1.1 开源数据	17
2.1.2 文本文件	18
2.1.3 Excel 文件	18
2.1.4 SQL 数据库	18
2.1.5 NoSQL 数据库	19
2.1.6 多媒体	20
2.1.7 网页检索	20

2.2	数据清洗	22	3.6.4	JavaScript	43
2.2.1	统计方法	23	3.6.5	SVG	43
2.2.2	文本解析	23	3.7	开始使用 D3.js	43
2.2.3	数据转化	25	3.7.1	柱状图	44
2.3	数据格式	25	3.7.2	饼图	48
2.3.1	CSV	26	3.7.3	散点图	50
2.3.2	JSON	27	3.7.4	单线图	52
2.3.3	XML	28	3.7.5	多线图	55
2.3.4	YAML	29	3.8	交互与动画	59
2.4	数据归约	30	3.9	社交网络中的数据	61
2.4.1	过滤及抽样	30	3.10	可视化分析的摘要	62
2.4.2	分箱算法	30	3.11	小结	62
2.4.3	降维	31			
2.5	开始使用 OpenRefine 工具	32	第 4 章	文本分类	63
2.5.1	text facet	33	4.1	学习和分类	63
2.5.2	聚类	33	4.2	贝叶斯分类	64
2.5.3	文本过滤器	34	4.3	E-mail 主题测试器	65
2.5.4	numeric facet	34	4.4	数据	66
2.5.5	数据转化	35	4.5	算法	68
2.5.6	数据输出	36	4.6	分类器的准确性	71
2.5.7	操作历史记录	36	4.7	小结	73
2.6	小结	37			
第 3 章	可视化	38	第 5 章	基于相似性的图像检索	74
3.1	可视化概述	39	5.1	图像相似性搜索	74
3.2	利用网页版的可视化	39	5.2	动态时间规整	75
3.3	探索科学可视化	39	5.3	处理图像数据集	77
3.4	在艺术上的可视化	40	5.4	执行 DTW	77
3.5	可视化生命周期	40	5.5	结果分析	79
3.6	可视化不同类型的数据	41	5.6	小结	81
3.6.1	HTML	41	第 6 章	模拟股票价格	82
3.6.2	DOM	42	6.1	金融时间序列	82
3.6.3	CSS	42	6.2	随机漫步模拟	83

6.3 蒙特卡罗方法	84	9.2 流行病模型	122
6.4 生成随机数	85	9.2.1 SIR 模型	122
6.5 用 D3.js 实现	86	9.2.2 使用 SciPy 来解决 SIR 模型的 常微分方程	123
6.6 计量分析师	91	9.2.3 SIRS 模型	124
6.7 小结	93	9.3 对细胞自动机进行建模	125
第 7 章 预测黄金价格	94	9.3.1 细胞、状态、网格和邻域	126
7.1 处理时间序列数据	94	9.3.2 整体随机访问模型	127
7.2 平滑时间序列	97	9.4 通过 D3.js 模拟 CA 中的 SIRS 模型	127
7.3 线性回归	100	9.5 小结	135
7.4 数据——历史黄金价格	101	第 10 章 应用社交图谱	136
7.5 非线性回归	101	10.1 图谱的结构	136
7.5.1 核岭回归	102	10.1.1 无向图	137
7.5.2 平滑黄金价格时间序列	104	10.1.2 有向图	137
7.5.3 平滑时间序列的预测	105	10.2 社交网络分析	137
7.5.4 对比预测值	106	10.3 捕获 Facebook 图谱	138
7.6 小结	107	10.4 使用 Gephi 再现图谱	139
第 8 章 使用支持向量机的方法 进行分析	108	10.5 统计分析	142
8.1 理解多变量数据集	109	10.6 度的分布	144
8.2 降维	111	10.6.1 图谱直方图	145
8.2.1 线性无差别分析	112	10.6.2 集中度	146
8.2.2 主成分分析	112	10.7 将 GDF 转化为 JSON	148
8.3 使用支持向量机	114	10.8 在 D3.js 环境下进行图谱可视化	150
8.3.1 核函数	115	10.9 小结	154
8.3.2 双螺旋问题	116	第 11 章 分析 Twitter 数据	155
8.3.3 在 mlpy 中实现 SVM	116	11.1 解析 Twitter 数据	155
8.4 小结	119	11.1.1 tweet	156
第 9 章 应用细胞自动机的方法 对传染病进行建模	120	11.1.2 粉丝	156
9.1 流行病学简介	120	11.1.3 热门话题	156

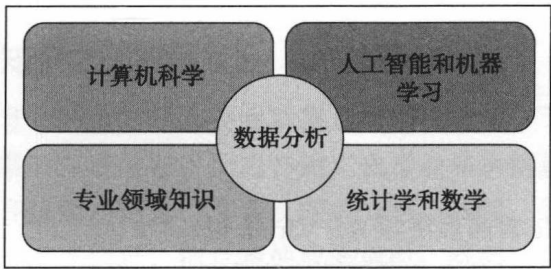
11.2 使用 OAuth 访问 Twitter API	157	13.3 在 MongoDB 中使用 MapReduce	190
11.3 开始使用 Twython	158	13.3.1 map 函数	190
11.3.1 利用 Twython 进行简单查询	159	13.3.2 reduce 函数	191
11.3.2 获取时间表数据	163	13.3.3 使用 Mongo shell	191
11.3.3 获取粉丝数据	165	13.3.4 使用 Jupyter	193
11.3.4 获取地点和趋势信息	167	13.3.5 使用 PyMongo	194
11.3.5 获取用户数据	168	13.4 过滤输入集合	195
11.3.6 API 流	169	13.5 分组和聚合	196
11.4 小结	171	13.6 在 tweet 中统计高频词汇	198
第 12 章 使用 MongoDB 进行数据		13.7 小结	201
处理和聚合	172	第 14 章 使用 Jupyter 和 Wakari	
12.1 开始使用 MongoDB	172	进行在线数据分析	202
12.1.1 数据库	173	14.1 开始使用 Wakari	202
12.1.2 集合	175	14.2 开始使用 Jupyter 记事本	205
12.1.3 文件	175	14.3 通过 PIL 进行图像处理	208
12.1.4 Mongo shell	175	14.3.1 打开图像	208
12.1.5 Insert/Update/Delete	176	14.3.2 显示图像直方图	208
12.1.6 查询	177	14.3.3 过滤	209
12.2 数据准备	178	14.3.4 操作	211
12.2.1 使用 OpenRefine 进行数据		14.3.5 转化	212
转换	179	14.4 开始使用 pandas	213
12.2.2 通过 PyMongo 插入文件	180	14.4.1 处理时间序列	213
12.3 分组	182	14.4.2 通过数据框架来操作多变量	
12.4 聚合框架	184	数据集	215
12.4.1 流水线	184	14.4.3 分组、聚合和相关	219
12.4.2 表达式	185	14.5 分享你的记事本	221
12.5 小结	186	14.6 小结	224
第 13 章 使用 MapReduce 方法	188	第 15 章 使用 Apache Spark 处理	
13.1 MapReduce 概述	188	数据	225
13.2 编程模型	189	15.1 数据处理平台	226

15.1.1 Cloudera 平台	226	15.3 Apache Spark 概述	231
15.1.2 安装 Cloudera VM	227	15.3.1 Spark 的生态系统	231
15.2 分布式文件系统概述	229	15.3.2 Spark 编程模型	232
15.2.1 使用 Hadoop 分布式文件 系统 (HDFS) 的具体步骤 ...	229	15.3.3 Apache 启动的介绍性操作 样例	234
15.2.2 利用 HUE 的 Web 界面来 进行文件管理	230	15.4 小结	235



开 始

数据分析是将原始数据进行排序和组织的过程，是用于帮助解释过去和预测未来的一系列方法。数据分析不只针对数字，而是关于如何设定或提出问题，演化解释，以及验证假设的过程。数据分析（Data Analysis）横跨了计算机科学、人工智能和机器学习、统计学和数学以及专业领域知识等多个领域，如下图所示。



这些技能都是很重要的，为了更好地了解问题及优化的解答，让我们来定义这些领域。

1.1 计算机科学

计算机科学为数据分析提供了工具。大规模数据的产生使计算机分析变得至关重要，对编程、数据库管理、网络管理和高性能计算的需求也逐步增多。有基本的 Python（或其他任何一种高级编程语言）编程经验对于理解本章的内容也是必不可少的。

1.2 人工智能

根据 Stuart Russell 和 Peter Norvig 二人的定义：

“人工智能（AI）一定与聪明的程序有关，那么让我们开始着手去编写这样的聪明程序吧。”

换句话说，人工智能研究那些可以模拟智能行为的算法。在数据分析中，我们应用人工智能来实施那些需要推理、相似性搜索或者无监督分类的智能活动。像是深度学习的领域就会依赖人工智能演算法，现在正使用的场景有聊天机器人、推荐引擎、图像分类，等等。

1.3 机器学习

机器学习（Machine Learning, ML）是对计算机的算法进行研究，从而使机器学会如何在特定情况或既定模式下进行反应。根据 Arther Samuel（1959）的定义：

“机器学习的研究领域是指在没有明确编程的情况下，赋予计算机进行学习的能力。”

机器学习所对应的一系列算法，根据该算法的训练过程，大致可以划分为三组：

- ❑ 有监督学习（supervised learning）
- ❑ 无监督学习（unsupervised learning）
- ❑ 强化学习（reinforcement learning）

这些算法贯穿本书，且伴随一些实际例子，引导读者了解从初始数据问题到编程化解决方案的全过程。

1.4 统计学

2009 年 1 月，Google 首席经济师 Hal Varian 说：

“我坚信未来十年最热门的职业将会是统计学家。人们以为我在开玩笑，但是谁会猜到在 20 世纪 90 年代计算机工程师曾是最热门的职业呢？”

统计学是对获取、分析和解读数据的方法加以发展和应用。数据分析包括了一系列统计学的技巧，例如：模拟、贝叶斯方法、预测、回归、线性分析及分类等。

1.5 数学

数据分析利用诸多种数学技术，如线性代数（向量和矩阵、因式分解以及特征值）、数值法和条件概念算法等。在本书中，每一章都是独立的并包含必要的数学知识。