



SPSS Modeler+Weka 数据挖掘从入门到实战

经管之家 主编 李御玺 唐绍祖 马伯 曾珂 编著

面向商业数据挖掘建模分析人员的教材
从具体的商业数据分析案例入手，轻松掌握数据挖掘的目的、方法、工具。

CDA数据分析师 系列丛书

SPSS Modeler+Weka 数据挖掘从入门到实战

经管之家 主编 李御玺 唐绍祖 马伯 曾珂 编著



电子工业出版社

Publishing House of Electronics Industry

北京·BEIJING

内 容 简 介

本书是一本面向商业数据挖掘建模分析人员的教材,从具体的商业数据分析案例入手,帮助读者掌握数据挖掘的目的、方法、工具与分析步骤。本书所采用的分析工具为目前颇受好评的IBM SPSS Modeler及开源软件Weka。IBM SPSS Modeler有很好的用户接口,也有不错的分析功能,但不具有前沿技术的分析模块,以及很难与现有的信息系统结合,而Weka恰能弥补其不足。同时,这两个软件都不需要编程,适合初学者。本书具体内容由四位活跃在数据挖掘教学和项目开发领域的人员完成,内容侧重软件的实际操作。本书力图以浅显的方式解释复杂的技术原理,尽量避免涉及过多的数学内容。

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

SPSS Modeler+Weka 数据挖掘从入门到实战 / 经管之家主编; 李御玺等编著. —北京: 电子工业出版社, 2019.5

(CDA 数据分析师系列丛书)

ISBN 978-7-121-31911-2

I. ①S… II. ①经… ②李… III. ①统计分析—应用软件 IV. ①C819

中国版本图书馆 CIP 数据核字 (2017) 第 137107 号

策划编辑: 张慧敏

责任编辑: 王 静

印 刷: 北京京师印务有限公司

装 订: 北京京师印务有限公司

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: 787×980 1/16 印张: 17.5 字数: 398 千字

版 次: 2019 年 5 月第 1 版

印 次: 2019 年 5 月第 1 次印刷

定 价: 69.00 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, 88258888。

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

本书咨询联系方式: 010-51260888-819, faq@phei.com.cn。

前言

感谢您选择《SPSS Modeler+Weka 数据挖掘从入门到实战》。本书内容源于李御玺教授的数据挖掘相关课程讲义，讲义历经多次修改，逐渐适合作为数据挖掘实用教材，并在获得学员们的高度评价后再被编辑成书。本书的另一位作者常国珍也长期活跃在数据挖掘的项目实施和培训中，2014 年其与李教授相识，并与李教授对出版本书之事一拍即合。

读者对象

本书的撰写采取了算法与软件实操双向并行的策略。在理论上，本书尽量用例子来说明数据挖掘算法背后的理论及意义，避免艰涩的数学公式，以求读者能用最简单的方式理解理论的精髓。在软件实操上，本书以各领域的实用案例为基础，逐步地将软件的功能引出，以求读者能了解软件功能的使用场景。有了坚实的理论基础及软件操作能力，再辅之以众多的实用案例，本书的读者就能逐步进入多姿多彩的数据挖掘世界。本书是以读者第一次接触数据挖掘为前提来撰写的。读者若有数据库、统计及计算机基础，则学习起来会较为轻松。

工具介绍

IBM SPSS Modeler 可谓商业数据挖掘领域的“重型武器”，其功能全面、算法安全可靠、追求执行效率与操作上的简单易用，并被广泛运用于许多企业中。其缺点是缺乏前沿的分析模块及很难与现有的信息系统结合，而开源软件 Weka 恰能弥补其不足。Weka 简单好用，拥有许多前沿的分析模块并易于与现有的信息系统整合。其缺点是在数据预处理部分，便利性不如 IBM SPSS Modeler 简单、易用。这两个软件对初入数据分析领域的读者而言是很好的入门工具。

阅读指南

本书分为 15 章。第 1 章介绍数据挖掘的起源及应用。同时说明如何建立一个 SPSS Modeler 及 Weka 的项目。第 2 章介绍数据挖掘的方法论 CRIPS-DM。同时说明如何将数据汇入 SPSS Modeler 及 Weka 的项目中，并做初步的数据探索。第 3 章介绍基本的数据挖掘技术。同时说明如何利用 SPSS

Modeler 及 Weka 建立 KNN 模型并进行分类预测。第 4 章介绍数据挖掘的进阶技术、数据挖掘技术的绩效增益及两个重要的数据挖掘网站。第 5 章详细介绍数据预处理技术，同时说明如何利用 SPSS Modeler，针对银行的信用风险评估数据，进行数据预处理。第 6 章介绍如何有效地挖掘对项目有帮助的关键变量。同时说明如何利用 SPSS Modeler 及 Weka，挖掘有效变量。第 7 至 15 章则为数据挖掘模型的介绍。这些模型均为热门且应用最为广泛的模型。对于每个模型的介绍，先以实例说明其理论，随后以实用的案例介绍如何在 SPSS Modeler 及 Weka 中操作，让每个读者充分了解每个模型的实际运用效果。

如果时间允许，则读者可以采取通读本书内容并按照示例进行操作的方式，但是这样效率可能不高。更高效的方法是结合工作中遇到的问题，先集中精力把书上的示例操练好，然后带入工作中的实际数据实现同样的算法，最后修改部分设置，以满足工作中的特定需求。

本书特点

本书作为市场上为数不多的理论与软件实操相结合并面向商业数据挖掘的书籍，和其他统计软件图书有很大的不同，本书结构新颖，案例贴近实际，讲解深入透彻。

- 场景式设置

本书从银行、电信、零售、医疗等行业中精心归纳、提炼出各类数据挖掘案例，方便读者搜寻与实际工作相似的问题。

- 启发式描述

本书注重培养读者解决问题的思路，以最朴实的思维方式结合启发式的描述，帮助读者发现规律、总结规律和运用规律，从而启发读者快速找出问题的解决方法。

售后服务

尽管作者们对书中的案例精益求精，但疏漏之处在所难免，如果发现书中的错误或某个案例有更好的解决方案，则敬请与本书作者联系，作者邮箱为 leeys@mail.mcu.edu.tw。

学习方法

只有对数据分析的流程熟悉了，才能实现从模仿到灵活运用。在产品质量管理方面，对流程的掌控是成功的关键，在数据挖掘项目中，流程同样是重中之重。数据挖掘是一个先后衔接的过程，一个步骤的失误会带来完全错误的结果。数据挖掘的流程大致包括抽样、数据清洗、数据转换、建模和模型评估这几个步骤。如果在抽样中的取数逻辑不正确，就有可能使因果关系倒置，得到完全

相反的结论。数据转换方法如果选择不正确，模型就难以得到预期的结果。而且，数据挖掘是一个反复试错的过程，每一步都要求有详细的记录和操作说明，否则分析人员很可能迷失方向。

学习数据挖掘最好的方法就是动手做一遍。本书语言通俗，但高度凝练，很少涉及公式，这会让读者大意，如果读者不动手做一遍，则很难体会到书中表述的思想。本书提供了相应的演练数据，也同时给出了相关方面的参考资料，供学员学习。

致谢

本丛书从策划到出版，张慧敏主编倾注了大量心血，经管之家的董事长赵坚毅先生提供了多方面的支持，特在此表示衷心的感谢！

为保证丛书的质量，使其更贴近读者，我们邀请了北京大学的殷子涵进行试读和修改完善。感谢各位预读员的辛勤、耐心与细致，使得本书能以更加完善的面目与各位读者见面。还要感谢刘莎莎参与本出的编写工作。

再次感谢您的支持！

作者

读者服务

轻松注册成为博文视点社区用户（www.broadview.com.cn），扫码直达本书页面。

- **提交勘误：**您对书中内容的修改意见可在 [提交勘误](#) 处提交，若被采纳，将获赠博文视点社区积分（在您购买电子书时，积分可用来抵扣相应金额）。
- **交流互动：**在页面下方 [读者评论](#) 处留下您的疑问或观点，与我们和其他读者一同学习交流。

页面入口：<http://www.broadview.com.cn/31911>

目 录

第 1 篇 理论篇

第 1 章 数据挖掘简介	1
1.1 数据挖掘的起源、定义及目标	2
1.2 数据挖掘的发展历程	2
1.3 SPSS Modeler 和 Weka 基础操作	4
1.3.1 SPSS Modeler 软件简介	4
1.3.2 建立一个 SPSS Modeler 项目	5
1.3.3 Weka 软件环境简介	8
1.3.4 Weka 简单操作实例	9
第 2 章 数据挖掘方法论	15
2.1 数据挖掘方法论	16
2.1.1 CRISP-DM	16
2.1.2 SEMMA	16
2.2 数据库中的知识挖掘步骤	17
2.2.1 字段选择	17
2.2.2 数据清洗	18
2.2.3 字段扩充	18
2.2.4 数据编码	19
2.2.5 数据挖掘	20
2.2.6 结果呈现	21
2.3 案例：运用 SPSS Modeler 和 Weka 做客户的信用风险评分模型	22
2.3.1 案例说明	22
2.3.2 案例实操	23
2.3.3 运用 SPSS Modeler 进行初步的数据挖掘	28

2.3.4	运用 Weka 进行数据汇入	34
2.3.5	Weka 自有数据存储格式 arff 简介	36
第 3 章	基本的数据挖掘技术	38
3.1	描述性统计	39
3.1.1	案例：通过数据判断客户是否需要新增电话线路	39
3.1.2	案例：运用描述性统计分析杂志社的客户特征	40
3.2	可视化技术	42
3.3	KNN 原理及实例	44
3.3.1	KNN (K 最近邻) 算法	44
3.3.2	使用 KNN 算法计算距离	45
3.3.3	案例：使用 KNN 算法向用户推荐电影	49
3.4	案例：运用 Weka 的 KNN 算法对诊断结果进行预测	52
3.4.1	案例说明	52
3.4.2	运用 Weka 中的 IBk 模型进行预测	53
3.5	案例：运用 SPSS Modeler 的 KNN 算法预测客户是否接受人寿保险推销	58
3.5.1	案例说明	58
3.5.2	案例实操	59
第 4 章	数据挖掘进阶技术	68
4.1	数据挖掘的功能分类	69
4.1.1	描述型数据挖掘 (无监督数据挖掘)	69
4.1.2	预测型数据挖掘 (有监督数据挖掘)	70
4.2	数据挖掘的绩效增益	72
4.2.1	数据挖掘模型评估指标：正确率、响应率、查全率、 F 值	72
4.2.2	数据挖掘模型评估指标：Gain Chart	74
4.2.3	数据挖掘模型评估指标：Lift Chart	75
4.2.4	数据挖掘模型评估指标：Profit Chart	76
4.3	数据挖掘网站	77
4.3.1	KDnuggets	77
4.3.2	Kaggle	80
4.4	案例：评估新产品的促销活动效果	82
4.4.1	案例说明	83
4.4.2	数据及字段描述	83
4.4.3	效能评估方式	85
4.4.4	比赛结果排名	85

第 2 篇 准备篇

第 5 章 数据预处理..... 87

5.1 字段选择..... 88

5.1.1 数据整合..... 88

5.1.2 数据过滤..... 88

5.1.3 案例：运用 SPSS Modeler 过滤数据..... 89

5.2 数据清洗..... 92

5.2.1 错误值的检测及处理..... 92

5.2.2 案例：运用 SPSS Modeler 进行错误值的检测及处理..... 92

5.2.3 离群值的检测及处理..... 96

5.2.4 案例：运用 SPSS Modeler 进行离群值的检测及处理..... 96

5.2.5 缺失值的检测及处理..... 100

5.2.6 案例：运用 SPSS Modeler 进行缺失值的检测及处理..... 101

5.3 字段扩充..... 110

5.3.1 案例说明..... 110

5.3.2 案例：运用 SPSS Modeler 进行字段扩充及评估对效能的提升..... 111

5.4 数据编码..... 118

5.4.1 数据转换..... 118

5.4.2 数据精简..... 128

5.4.3 数据集的切割..... 129

第 6 章 关键变量挖掘技术..... 137

6.1 无效变量..... 138

6.2 统计方式的变量选择..... 138

6.2.1 卡方检验..... 138

6.2.2 方差分析（ANOVA 检验）及 *t* 检验..... 138

6.2.3 案例：运用 SPSS Modeler 进行关键变量挖掘..... 139

6.3 模型方式的变量选择..... 141

6.3.1 决策树..... 141

6.3.2 Logistic 回归..... 141

第 7 章 贝叶斯网络..... 143

7.1 朴素贝叶斯..... 144

7.1.1 独立性假设	145
7.1.2 概率的离散化	147
7.2 什么是贝叶斯网络	147
第 8 章 线性回归	150
8.1 简单线性回归	151
8.2 多元回归	152
8.3 相关系数	152
8.4 回归分析案例	153
8.5 线性回归模型评估	156
8.5.1 线性回归模型评估指标: MAE、MSE 和 RMSE	156
8.5.2 线性回归模型评估指标: R^2	156
8.6 案例: 运用 SPSS Modeler 建立线性回归模型	157
8.6.1 案例说明	157
8.6.2 案例实操	157
第 9 章 决策树	161
9.1 ID3 决策树模型	162
9.2 ID3 算法	165
9.2.1 ID3 算法的字段选择方式	165
9.2.2 使用决策树进行分类	168
9.2.3 决策树与决策规则之间的关系	168
9.2.4 ID3 算法的缺点	169
9.3 C5.0 算法	170
9.3.1 C5.0 算法的字段选择方式	170
9.3.2 C5.0 算法的数值型字段处理方式	170
9.3.3 C5.0 算法的剪枝方法	172
9.4 CART 算法	173
9.4.1 分类树与回归树	174
9.4.2 CART 分类树的字段选择方式	174
9.4.3 CART 分类树的剪枝作法	177
9.5 CHAID 算法	177
9.6 案例: 运用 SPSS Modeler 和 Weka 建立决策树模型	177
9.6.1 案例说明	177
9.6.2 案例实操	178

9.6.3	运用 SPSS Modeler 建立交互式分类树模型	179
9.6.4	运用 Weka 建立交互式分类树模型	180
9.7	CART 回归树算法	186
9.7.1	CART 回归树的字段选择方式	186
9.7.2	利用模型树提升 CART 回归树的效率	187
9.8	案例：运用 SPSS Modeler 和 Weka 建立回归树模型	188
9.8.1	案例说明	188
9.8.2	案例实操	188
9.8.3	使用 Weka 对比“剪枝”前后的模型	189
第 10 章 神经网络		194
10.1	BP 神经网络模型	195
10.1.1	BP 神经网络模型的概念	195
10.1.2	BP 神经网络模型的架构方式	195
10.2	神经元的组成	198
10.3	神经网络模型如何传递信息	199
10.4	修正神经网络模型的权重值及常数项	200
10.5	BP 神经网络模型与 Logistic 回归、线性回归及非线性回归之间的关系	201
10.6	案例：运用 SPSS Modeler 建立类神经网络模型	202
第 11 章 Logistic 回归		208
11.1	Logistic 回归与 BP 神经网络的关系	210
11.2	Logistic 回归的字段选择方式	211
11.2.1	前向法	211
11.2.2	后向法	212
11.2.3	逐步法	212
11.3	案例：运用 SPSS Modeler 建立 Logistic 回归模型	213
11.3.1	案例说明	213
11.3.2	案例实操	213
第 12 章 支持向量机		215
12.1	数据是线性可分的支持向量机	217
12.2	数据是线性不可分的支持向量机	219
12.3	案例：运用 SPSS Modeler 建立 SVM 模型	221

第 3 篇 关系篇

第 13 章	聚类分析	230
13.1	相似性度量	232
13.1.1	二元变量的相似性度量	232
13.1.2	类别型变量的相似性度量	234
13.1.3	数值型变量的相似性度量	234
13.2	聚类算法	234
13.2.1	互斥聚类与非互斥聚类算法	234
13.2.2	分层聚类算法	235
13.2.3	分割式聚类算法	236
13.3	分层聚类算法	236
13.3.1	单一连接法	236
13.3.2	完全连接法	237
13.3.3	平均连接法	238
13.3.4	中心法	238
13.3.5	Ward's 法 (华德法)	239
13.4	分割式聚类算法	240
13.4.1	K-Means 算法	240
13.4.2	K-Medoids 算法	243
13.4.3	SOM 算法	243
13.4.4	两步法	243
13.5	集群判断	244
13.5.1	集群判断方法: R^2	244
13.5.2	集群判断方法: 半径 R^2	245
13.5.3	集群判断方法: 均方根标准差 (RMSSTD)	245
13.6	案例: 运用 SPSS Modeler 建立聚类模型	246
13.6.1	案例说明	246
13.6.2	案例实操	246
第 14 章	关联规则	252
14.1	关联规则的概念	253
14.2	关联规则的评估指标	253
14.2.1	支持度	253
14.2.2	置信度	254

14.3	Apriori 算法	254
14.3.1	暴力法的问题	254
14.3.2	Apriori 算法的理论基础	255
14.4	Apriori 算法实例说明	255
14.4.1	候选项目组合的产生	255
14.4.2	候选项目组合的删除	256
14.5	再谈评估指标	256
14.5.1	支持度与置信度的问题	256
14.5.2	提升度指标	257
14.6	关联规则的延伸	257
14.6.1	虚拟商品的加入	257
14.6.2	负向关联规则	257
14.7	案例：运用 SPSS Modeler 建立关联规则模型	258
14.7.1	案例说明	258
14.7.2	案例实操	258
 第 15 章 序列模型		263
15.1	序列模型的概念	264
15.2	案例：运用 SPSS Modeler 建立序列模型	266
15.2.1	案例说明	266
15.2.2	案例实操	266

第 1 篇 理论篇

第 1 章

数据挖掘简介

近年来，信息产业高速发展，人们越来越关注如何将信息转换成有用的、直观的知识。因此，在 1991 年，William Frawley 和 Gregory Piatetsky Shapiro 提出了数据挖掘的概念，即从现有的大量数据中，撷取不明显的、之前未知的、可能有用的知识的过程。数据挖掘的目标是建立一个决策模型，根据过往的行动来预测未来的行为。例如，分析一家公司的不同客户对公司产品的购买情况，进而分析出哪一类客户会对公司的产品有兴趣。在讲究实时、竞争激烈的网络时代，若能事先破解消费者的行为模式，将是公司获利的关键因素之一。

1.1 数据挖掘的起源、定义及目标

数据挖掘可以从数据中撷取不明显的、之前未知的信息，举一个著名的例子：美国的沃尔玛超市为了能够准确了解顾客在其门店的购买习惯，对其顾客的购物行为进行数据挖掘。一个意外的发现是：跟尿布一起销售最多的商品竟是啤酒！在美国，一些年轻的父亲下班后经常要到超市去买婴儿尿布，而他们中有 30%~40%的人同时也为自己买一些啤酒。产生这一现象的原因是：美国的太太们常叮嘱她们的丈夫在下班后为小孩买尿布，而丈夫们在买尿布后又随手带回了他们喜欢的啤酒。既然尿布与啤酒被一起购买的机会很多，于是，沃尔玛就在其一个个门店中将尿布与啤酒并排摆放在一起，结果是尿布与啤酒的销售量双双增长。按照常规思维，尿布与啤酒风马牛不相及，若不是借助数据挖掘技术对大量交易数据进行分析，沃尔玛是不可能发现这个有价值的规律的。

与此同时，数据挖掘能够发现信息之间的联系，但不一定是因果关系，在数据挖掘的过程中能发现有用的信息。比如，美国的分析员做过这样一项研究：在冰激凌热卖时，溺水的人数会增加。而冰激凌与溺水的人数之间存在的联系并不是因果关系，所以，减少冰激凌的销售量不一定会降低溺水的人数。而二者之所以会存在联系，是因为冰激凌在夏天热卖，而夏天去游泳的人数增加，溺水的人数也会增加。所以，通过这个例子可以看出，数据挖掘要在大量的数据中找到有用的知识，不能发现关联之后就随意下定结论。

1.2 数据挖掘的发展历程

数据挖掘虽然是计算机应用领域的新名词，但也经历了几十年的发展历程。

- 第一阶段：1960 年以后，各种新兴的数据收集模式开始出现，例如磁带、软盘、硬盘等，人们开始掌握了收集数据的基本方法。
- 第二阶段：1980 年以后，随着收集的数据量的增多，人们开始需要数据库，并逐渐建立起了数据库，但是此时还不能查询数据。
- 第三阶段：1990 年以后，数据统计的概念出现，人们可以进入数据仓库完成简单的数据统计，但并不能做太精细的决策。
- 第四阶段：2000 年以后，随着数据库和计算机网络的广泛应用，加上使用先进的自动数据生成和采集工具，人们所拥有的数据量急剧增大。针对大规模数据的分析处理方法——数据挖掘出现了。

数据挖掘在各行各业中都有应用，比如其最早应用于银行、通信业，现在也在零售业、保险业及政府中有所应用（见图 1-1）。

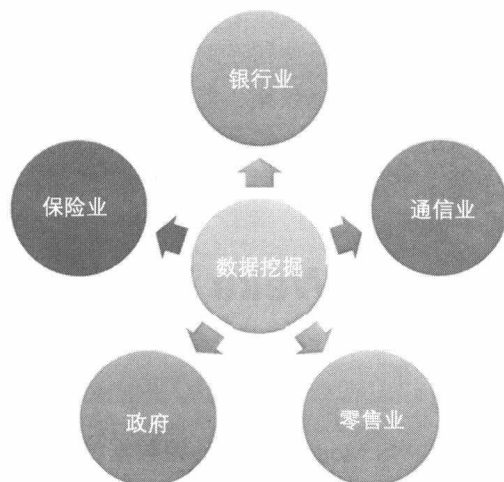


图 1-1

1. 银行

金融事务需要收集和处理大量数据，由于银行在金融领域的地位、工作性质及业务特点，市场竞争激烈程度决定了它对信息化、电子化的需求比其他领域更迫切。利用数据挖掘技术可以帮助银行产品开发部门描述客户以往的需求并预测未来。例如，汇丰银行对不断增长的客户群进行分类，为每种产品找到最有价值的客户，这样其产品才能推销得好，而且比盲目推销产品节省了 30% 的销售费用。再例如，银行通过数据挖掘发现，有盗刷信用卡行为的人，其使用信用卡的第一笔消费往往小于 10 元，所以，银行根据这条规律，冻结了那些第一笔消费小于 10 元的客户的账号，减少了客户盗刷信用卡所带来的经济损失。

2. 零售业

在过去，零售商依靠供应链软件、内部分析软件甚至直觉来预测库存需求。随着竞争压力一天天地增大，很多零售商都开始致力于找到更准确的方法来预测其连锁商店应保有的库存。通过数据挖掘可以为产品存储决策提供准确、及时的信息。

3. 保险业

对受险人员进行分类有助于确定适当的保险金额度。通过数据挖掘可以得到不同行业、不同年龄段、不同社会层次的人的信息，从而可以评估他们的保险金。另外，还可进行保险金种类关联分析，分析购买了某种保险的人是否又同时购买了另一种保险，也可预测什么样的顾客会购买新险种。总而言之，数据挖掘在保险业中有广泛的应用。

4. 政府

数据挖掘被广泛应用于电子政务中的综合查询、经济分析、宏观预测、应急预警、风险分析及预警、质量监督管理及监测、决策支持等系统，它为公众提供了一个智能化、高效的网上政府。例

如，几乎每个政府网站都有类似“公众意见调查”的栏目，这是了解公众需求的一个很好的途径，但是从网站公布的调查结果看，其结论大多还停留在对单个问题求总数、求比例等简单分析上。利用数据挖掘技术可以在网上建立一个能有效地收集、监测和分析公众数据的系统，提炼出实用、有效的信息，建成以公众需求为主导的电子政务。将数据挖掘技术引入电子政务，可大大提高整个电子政务系统的智能化水平。

1.3 SPSS Modeler 和 Weka 基础操作

1.3.1 SPSS Modeler 软件简介

IBM SPSS Modeler（以下简称 SPSS Modeler）是一组数据挖掘工具，通过它可以快速建立预测模型，并将其应用于商业活动中，从而改进决策过程。

使用 SPSS Modeler 处理数据主要分为 3 个步骤。

- 首先，将数据读入 SPSS Modeler 中。
- 然后，通过一系列操作运行数据。
- 最后，将数据发送到目标位置。

这个操作过程被称为数据流，因为数据以一条条记录的形式，依次经过各种操作，最终到达目标位置（模型或某种数据输出）。

数据流工作区是 SPSS Modeler 窗口中最大的区域，也是构建和操作数据流的区域，如图 1-2 所示。

SPSS Modeler 中的大部分数据和建模工具位于节点选项卡中，该选项卡位于数据流工作区（简称工作区）的底部（见图 1-3）。要将节点添加到工作区中，在节点选项卡中双击节点对应的图标或将其拖曳到工作区中即可。随后可将各个图标连接以创建一个数据流。每个选项卡中均包含一组不同的数据流操作阶段中使用的相关节点，例如：

- 源（Source）节点。此类节点可将数据引入 SPSS Modeler 中。
- 记录选项（Record Ops）节点。此类节点可对数据记录执行操作，例如选择、合并和追加等。
- 字段选项（Field Ops）节点。此类节点可对数据字段执行操作，例如过滤、导出新字段和确定给定字段的测量级别等。
- 图形（Graphs）节点。此类节点可在建模前后以图表形式显示数据。图表形式包括散点图、直方图、网络节点和评估图表。
- 建模（Modeling）节点。此类节点可使用 SPSS Modeler 中提供的建模算法，例如神经网络、决策树、聚类算法和数据排序等。
- 输出（Output）节点。此类节点可生成能在 SPSS Modeler 中查看的数据、图表和模型等多种输出结果。