

数据统计分析软件SPSS入门

四川大学图书馆
舒予

主要内容

□1-SPSS概述

□2-SPSS数据管理

□3-SPSS的统计分析功能

主要内容

□1-SPSS概述

□2-SPSS数据管理

□3-SPSS的统计分析功能

1-SPSS概述

□1.1-SPSS简介

□1.2-窗口介绍

□1.3-数据库的构建

1-SPSS概述

□1.1-SPSS简介

□1.2-窗口介绍

□1.3-数据库的构建

1.1-SPSS简介

- SPSS是世界最早的统计分析软件，广泛应用于通信、医疗、银行、证券、保险、制造、商业、市场研究、科研、教育等许多领域和行业。
- SPSS的基本功能包括数据管理、统计分析、图表分析、输出管理等，具体的内容包括描述统计、总体的均值比较、相关分析、回归模型分析、聚类分析、时间序列分析、非参数检验等多个大类。

1.1-SPSS简介

□SPSS的特点

- 包括了各种成熟的统计方法与模型，为统计分析用户提供了全方位的统计学算法
- 提供了各种数据准备与数据整理技术
- 自由灵活的表格功能
- 各种常用的统计学图形
- 界面友好，操作简单，容易上手

1.2-窗口介绍

窗口启动



1.2-窗口介绍

□数据窗口

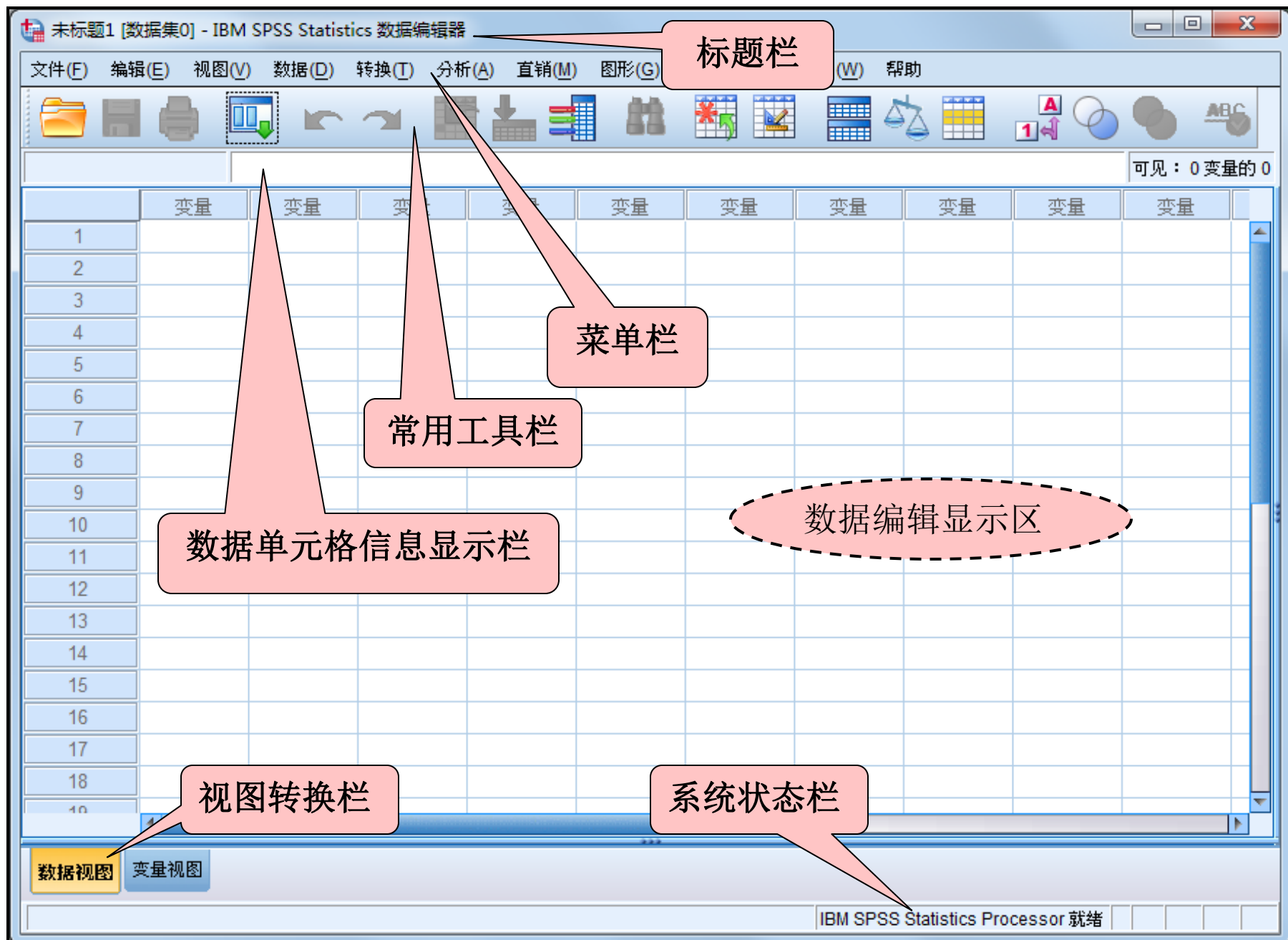
□变量窗口

□结果输出窗口

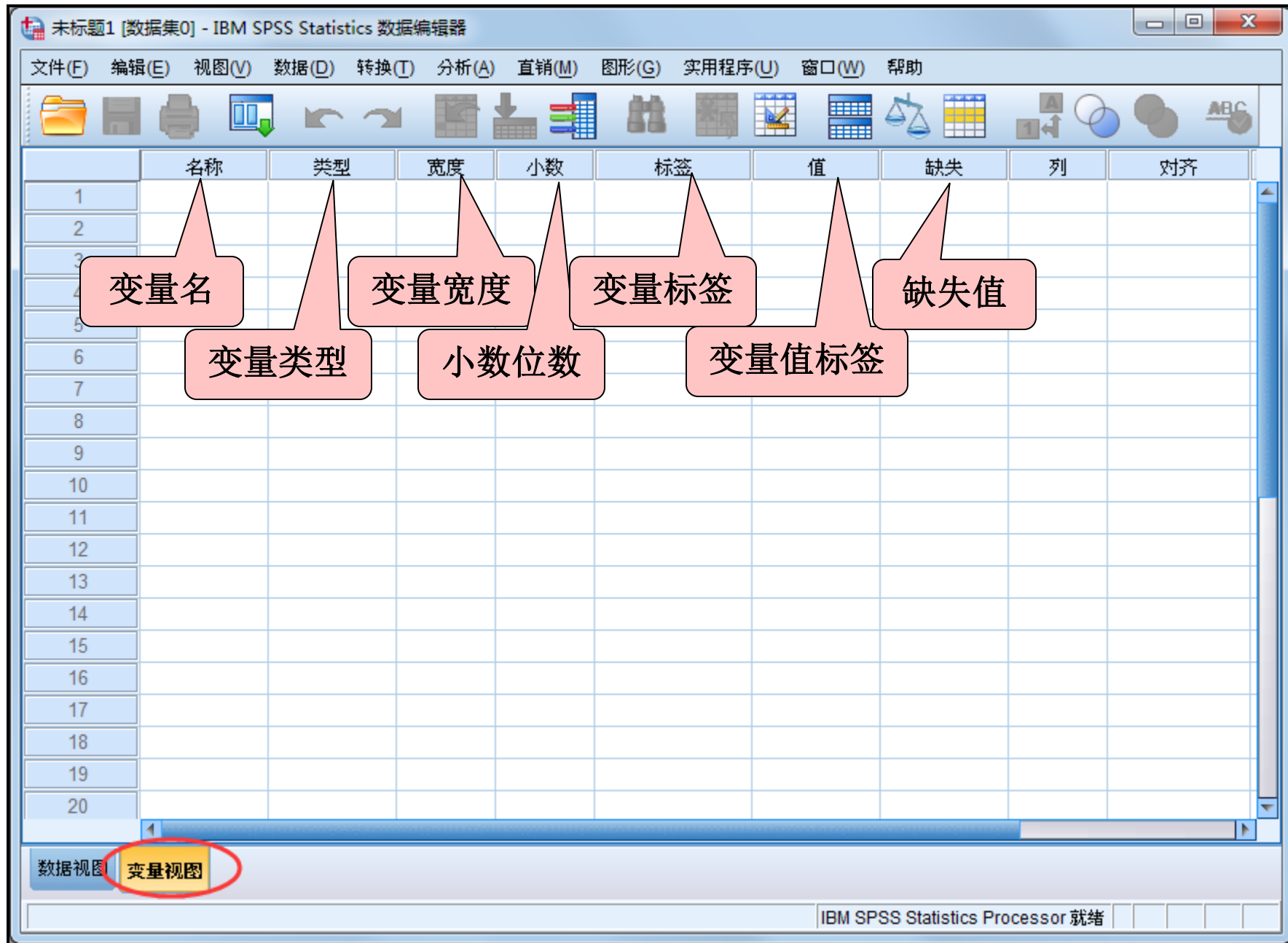
□图表编辑窗口

□程序编辑窗口

数据窗口



变量窗口



结果输出窗口

输出

- 日志
- T检验
 - 标题
 - 附注
 - 活动的数据集
 - 单个样本统计量
 - 单个样本检验

GET
FILE='E:\数据\2.sav'.
DATASET NAME 数据集1 WINDOW=FRONT.
T-TEST
/TESTVAL=0
/MISSING=ANALYSIS
/VARIABLES=浓度
/CRITERIA=CI(.95).

→ **T检验**

[数据集1] E:\数据\2.sav

单个样本统计量

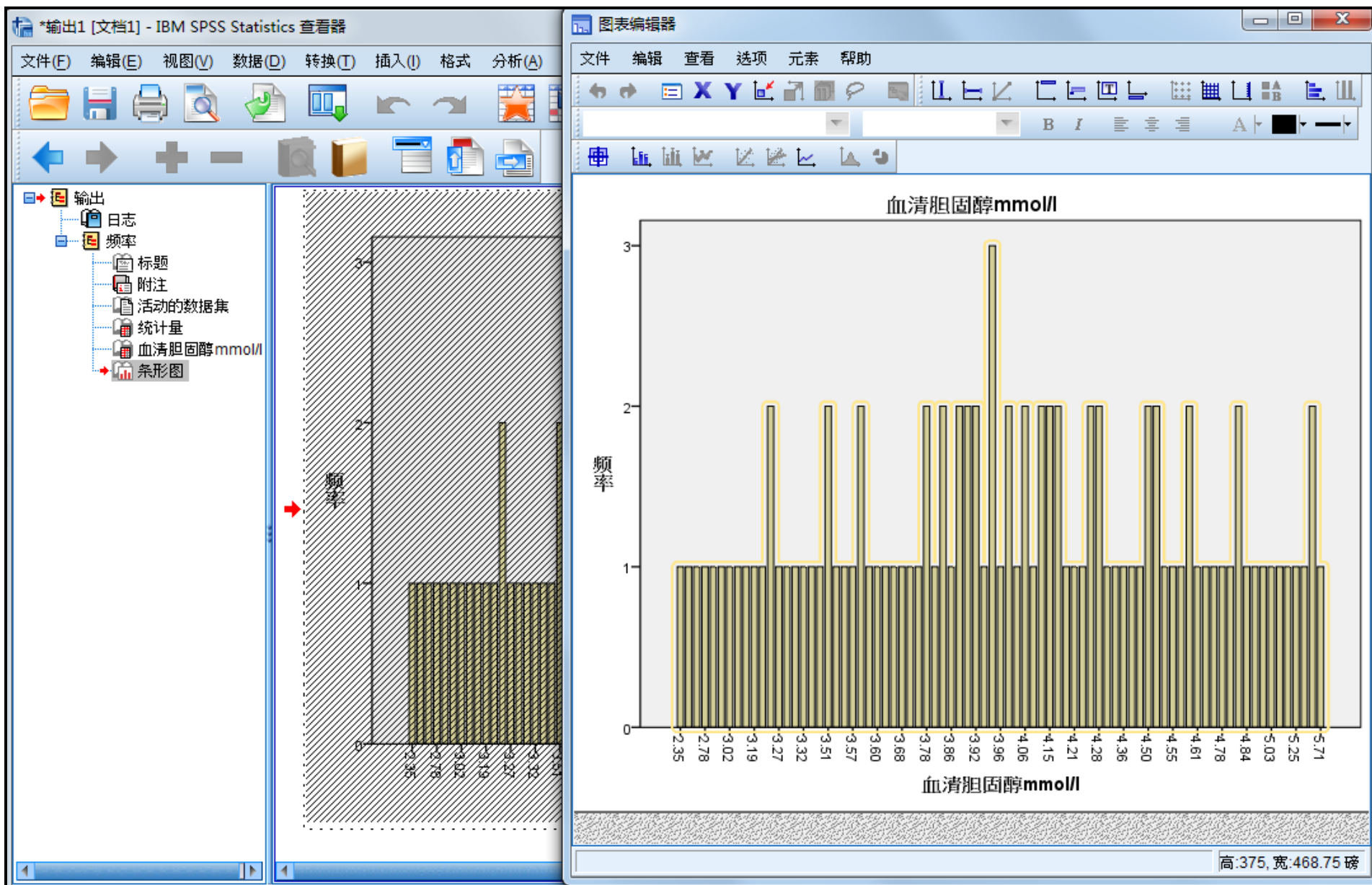
	N	均值	标准差	均值的标准误
浓度	11	20.9836	1.06750	.32186

单个样本检验

	检验值 = 0					
	t	df	Sig.(双侧)	均值差值	差分的 95% 置信区间	
					下限	上限
浓度	65.194	10	.000	20.98364	20.2665	21.7008

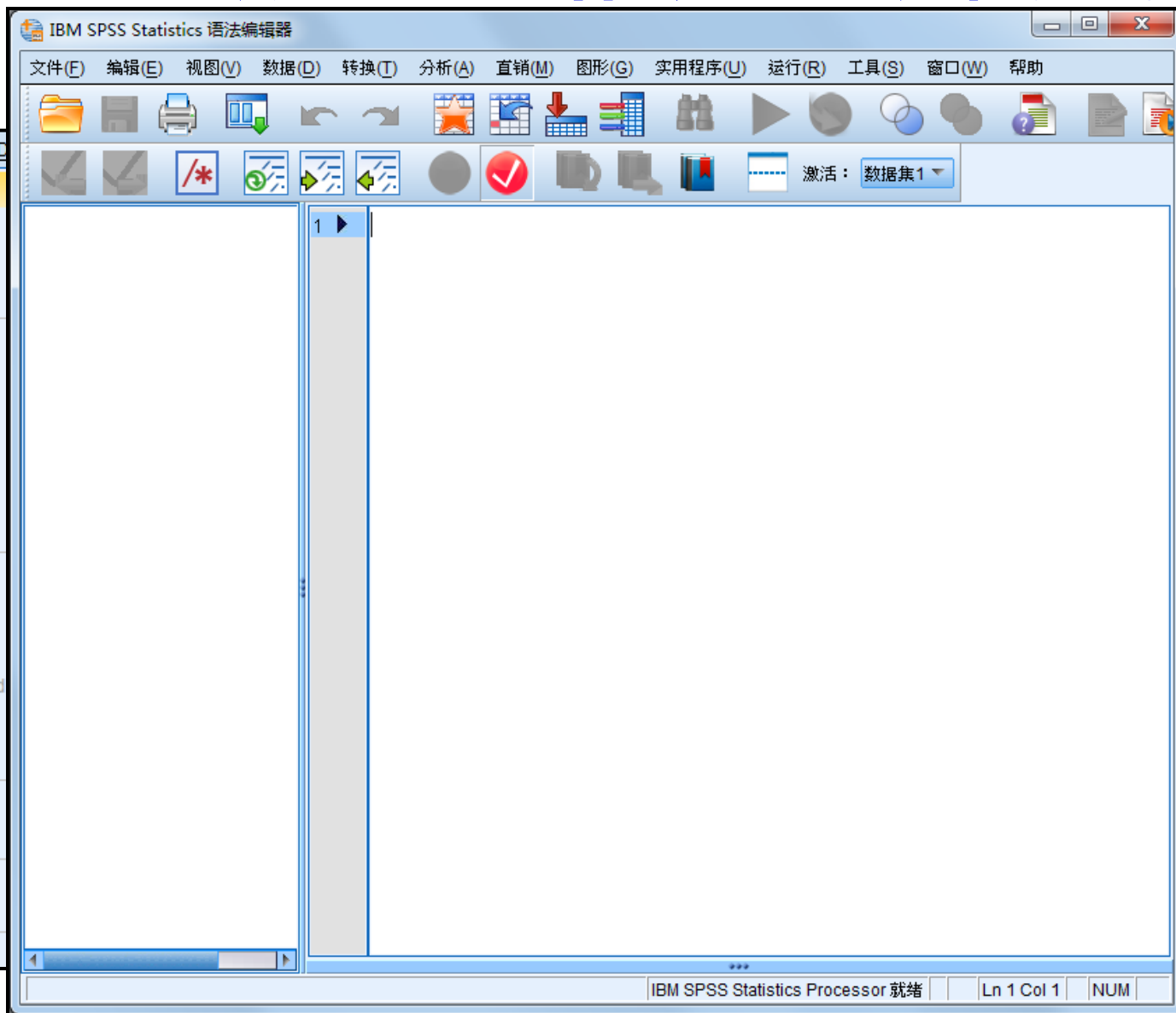
IBM SPSS Statistics Processor 就绪

图表编辑窗口



程序编辑窗口

□ “文件” — “新建” — “语法” 可打开程序编辑窗口



1.3-数据库的构建

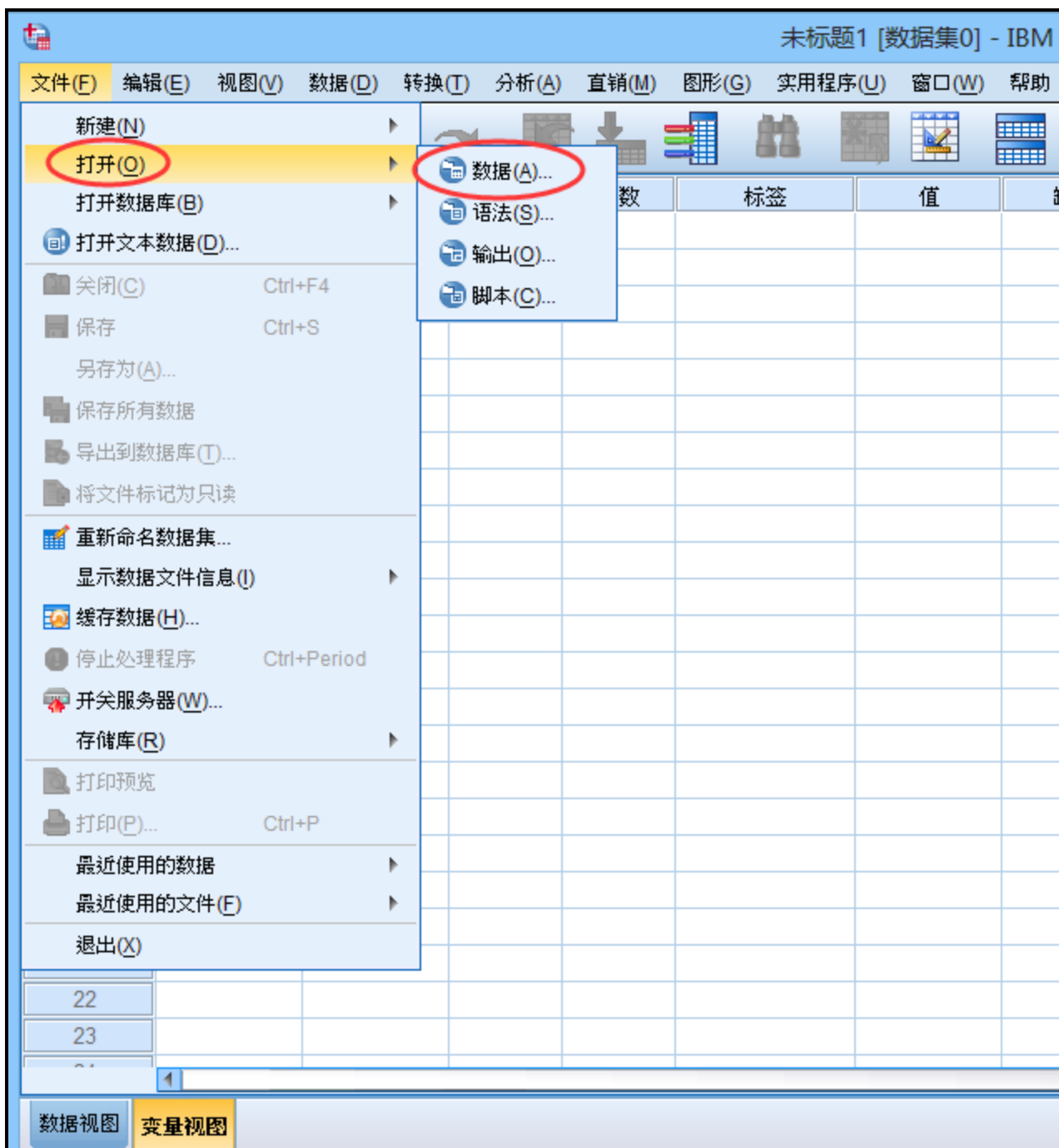
□现有文件导入

□直接录入数据

现有文件导入

□直接导入SPSS或其它形式的数据库

- 例1：导入例1. sav、例1. xls和例1. txt



直接录入数据

□例2：按照例2.xls的内容，构建SPSS数据库，并且性别记为男-1，女-0

直接录入数据

The image shows the IBM SPSS Statistics interface. The 'Data' menu is open, and the 'Data' option is highlighted. A red arrow points from the 'Data' menu to the 'Variable View' tab in the '未标题2 [数据集1] - IBM SPSS Statistics 数据编辑器' window. A text box with the text '首先对变量进行定义' (First, define the variables) points to the 'Variable View' tab.

未标题1 [数据集0] - IBM SPSS Statistics 数据编辑器

文件(F) 编辑(E) 视图(V) 数据(D) 转换(T) 分析(A) 直销(M) 图形(G) 实用程序(U) 窗口(W) 帮助

新建(N) 打开(O) 打开数据库(B) 打开文本数据(D)... 关闭(C) Ctrl+F4 保存 Ctrl+S 另存为(A)... 保存所有数据 导出到数据库(T)... 将文件标记为只读 重新命名数据集... 显示数据文件信息(I) 缓存数据(H)... 停止处理程序 Ctrl+Period 开关服务器(W)... 存储库(R) 打印预览 打印(P)... Ctrl+P 最近使用的数据 最近使用的文件(F) 退出(X)

未标题2 [数据集1] - IBM SPSS Statistics 数据编辑器

文件(F) 编辑(E) 视图(V) 数据(D) 转换(T) 分析(A) 直销(M) 图形(G) 实用程序(U) 窗口(W) 帮助

	名称	类型	宽度	小数	标签	值	缺失	列	对齐
1									
2									
3									
4									
5									
6									
7									
8									
9									
10									
11									
12									
13									
14									
15									
16									
17									
18									
19									
20									
21									
22									
23									

数据视图 变量视图

IBM SPSS Statistics Processor 就绪

变量类型和度量标准

 变量类型 ✕

☒ 数值(N)
☐ 逗号(C)
☐ 点(D)
☐ 科学计数法(S)
☐ 日期(A)
☐ 美元(L)
☐ 设定货币(U)
☐ 字符串(R)

宽度(W): 8
小数位(P): 2

确定 取消 帮助

度量标准

 度量(S) ▼

 度量(S)

 序号(O)

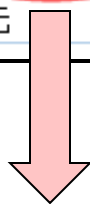
 名义(N)

□在SPSS中，每一个变量都有一个度量标准，这些度量标准说明变量的含义和属性，会对后续的分析产生影响。

- **度量：**度量表示定距变量和定比变量，这两类变量可以明确的表示事物之间的差值，拥有非常多的数据线信息，也是高级的测量水平，在统计分析中主要参与加减乘除的算术运算，其数据内容往往是数值。
- **序号：**序号表示定序变量，定序变量表示事物的顺序或等级，可以排序或比较优劣，可以计算频数和累计频率，定序变量的数据可以是数值，也可以是字符。
- **名义：**名义表示定类变量，定类变量表示事物的类别，只能计算频数和频率，各类别之间没有大小、顺序、等级之分。定类变量的数据可以是数值，也可以是字符。

对变量进行定义

名称	类型	宽度	小数	标签	值	缺失	列	对齐	度量标准
name	字符串	8	0	姓名	无	无	8	左	名义(N)
address	字符串	20	0	地址	无	无	8	左	名义(N)
age	数值(N)	8	0	年龄	无	无	8	右	度量(S)
sex	数值(N)	8	0	性别	{0, 女}...	无	8	右	度量(S)
score	数值(N)	8	1	分数	无	无	8	右	度量(S)



值标签(V)

值标签(V)

值(U):

拼写(S)...

标签(L):

添加(A)

更改(C)

删除(R)

0 = "女"

1 = "男"

确定 取消 帮助

□从Excel中直接复制数据进数据编辑区域

name	address	age	sex	score
张三	成都市一环路南一段24号	20	0	90.0
李四	成都市望江路29号	19	0	95.0
王五	成都市人民南路三段17号	21	0	80.0
赵六	成都市双流县川大路	20	0	85.0

数据库中对性别的定义是**数值型**，但是文件中的数据为**字符型**

解决方法1：

- 在变量视图中，将“性别”的变量类型改为“字符串”（与Excel文件中一致），并相应对值标签进行修改。再在数据视图选择“值标签”。

sex	字符串	8	0	性别	{女, 0}...
-----	-----	---	---	----	-----------

值标签(V)

值标签(V)

值(U):

标签(L):

女 = "0"
男 = "1"

添加(A)

编辑(E) 视图(V) 数据(D) 转换(T) 分析(A) 直销(M) 图形(G) 实用程序(U) 窗口(W) 帮助						
						
						
name	address	age	sex	score	变量	3
张三	成都市一环路南一段24号	20	1	90.0		
李四	成都市望江路29号	19	0	95.0		
王五	成都市人民南路三段17号	21	0	80.0		
赵六	成都市双流县川大路	20	1	85.0		

□解决方法2：

- 在Excel文件中，在性别一列中将“男”替换为“1”，“女”替换为“0”。
- 替换完成后，再复制进SPSS数据窗口。

□其它

- 变量位置的移动
- 添加、删除变量
- 和Excel的使用方式一致

主要内容

□1-SPSS概述

□2-SPSS数据管理

□3-SPSS的统计分析功能

2-SPSS数据管理

□2.1-数据的整理

2.1-数据的整理

□数据合并

□数据拆分

□数据排序

数据合并

□将若干小的数据文件合并成一个大的数据文件

□纵向合并

- 几个数据集中的数据纵向堆叠，组成一个新的数据集，新的数据集中的记录数是原来的几个数据集中记录数的总和
- 合并条件：（1）待合并的SPSS数据文件，其内容合并是有实际意义的；（2）不同数据文件中，数据含义相同的列，最好起相同的名字，变量类型和变量长度也尽量相同

数据合并

□横向合并

- 按照记录的次序，或者某个关键变量的数值，将不同数据集中不同变量合并为一个数据集，新数据集的变量数是所有原数据集中不重名变量的总和，实质就是将两个数据文件的记录，按照记录对应，一一进行左右对接，合并的两个数据文件的变量不同，但具有相同个案例数。
- 合并条件：（1）如果不是按照记录号对应的规则进行合并，则两个数据文件必须至少有一个变量名相同的公共变量，这个变量是这两个数据文件横向合并的依据，称为关键变量。（2）如果是使用关键变量进行合并的对应，则两个数据文件都必须事先按关键变量进行升序排列。（3）为方便SPSS数据文件的合并，在不同数据文件中，数据含义不相同的列，变量名不应取相同的名称。

数据合并

□例3：将例3-1.sav与例3-2.sav数据进行合并（纵向合并）

The screenshot illustrates the steps to merge data files in IBM SPSS. The main window, titled '例3-1.sav [数据集5] - IBM SPSS', shows the 'Data' menu circled in red. A red arrow points from this menu to a secondary window titled '添加个案(C)...' (Add Cases), which also has a red circle around its title bar. This secondary window contains two options: '添加个案(C)...' (Add Cases) and '添加变量(V)...' (Add Variables). The 'Data' menu is open, displaying various options, with '合并文件(G)' (Merge Files) circled in red. A red arrow points from '合并文件(G)' to the '添加个案(C)...' window. The 'Data' menu options include: 定义变量属性(V)..., 设置未知测量级别(L)..., 复制数据属性(C)..., 新建设定属性(B)..., 定义日期(E)..., 定义多重响应集(M)..., 验证(L), 标识重复个案(U)..., 标识异常个案(I)..., 排序个案..., 排列变量..., 转置(N)..., 合并文件(G), 重组(R)..., 分类汇总(A)..., 正交设计(H), 复制数据集(D), 拆分文件(F)..., 选择个案..., and 加权个案(W).... The data table in the background has columns '职工号' (Employee ID), '职称' (Title), and '工资' (Salary), with 12 rows of data.

	职工号	职称	工资
1	4	3	1246.00
2	1	1	1256.00
3	2	1	1326.00
4	3	2	1854.00
5	6	2	2422.00
6	7	3	2522.00
7	8	2	2552.00
8	9	1	2555.00
9	5	3	5624.00
10			
11			
12			

数据合并

例4：将例4-1.sav与例4-2.sav数据进行合并（横向合并）

例4-1.sav [数据集1] - IBM SPSS Statistics 数据编辑器

文件(F) 编辑(E) 视图(V) **数据(D)** 转换(T) 分析(A) 直销(M) 图形(G) 实用程序(U) 窗口(W) 帮助

定义变量属性(V)...
设置未知测量级别(L)...
复制数据属性(C)...
新建设定属性(B)...
定义日期(E)...
定义多重响应集(M)...
验证(L)
标识重复个案(U)...
标识异常个案(I)...
排序个案...
排列变量...
转置(N)...
合并文件(G)
重组(R)...
分类汇总(A)...
正交设计(H)
复制数据集(D)
拆分文件(F)...
选择个案...
加权个案(W)...

添加个案(C)...
添加变量(V)

可见：3 变量的 3

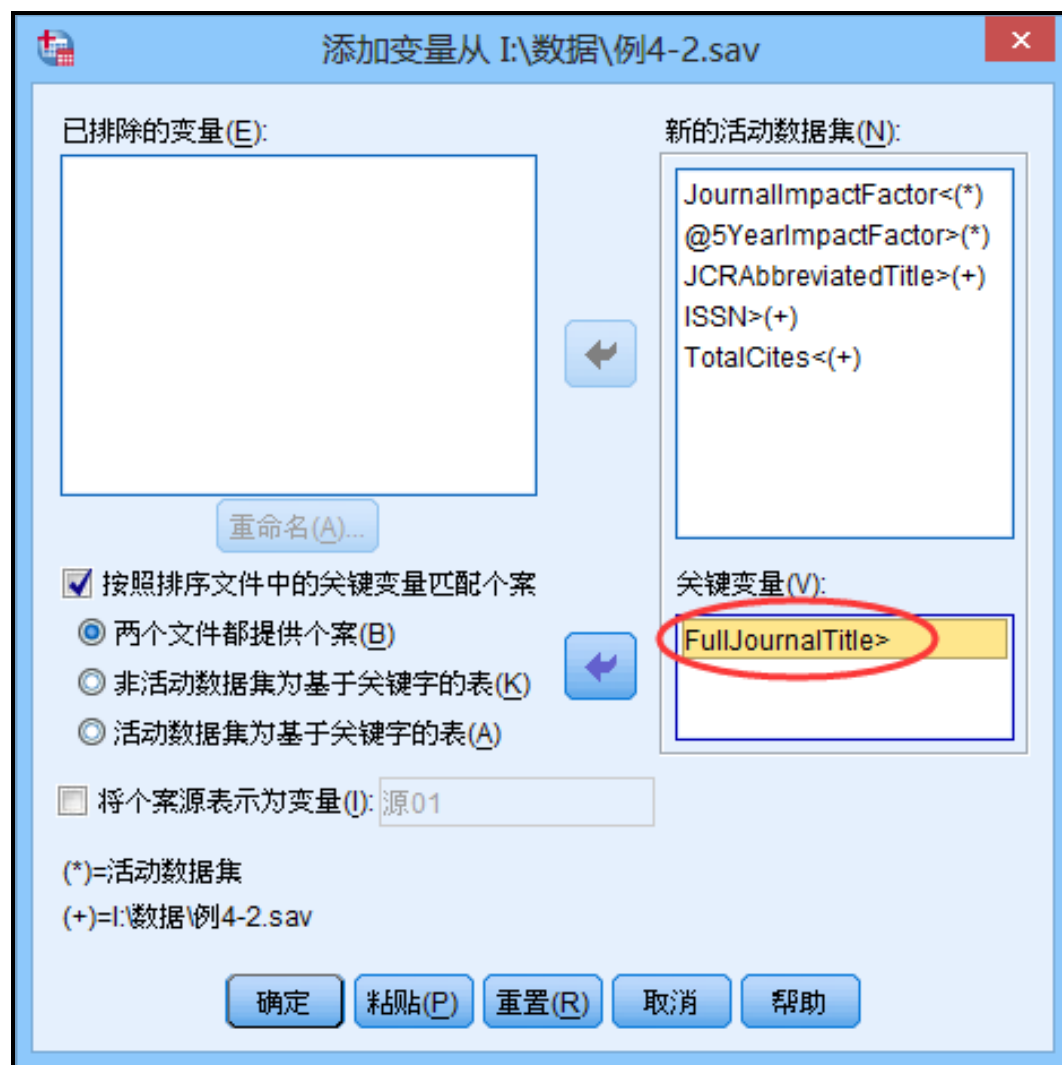
		Journal mpactFa ctor	@5YearImpactFactor	变量	变量	变量	变
1	TRAC-TREND	8.442	7.929				
2	BIOSENSORS	7.780	6.862				
3	Annual Review	7.435	9.259				
4	ANALYTICAL	6.320	6.016				
5	SEPARATION	6.077	5.327				
6	SENSORS AN	5.401	4.855				
7	ANALYTICA C	4.950	4.849				
8	MICROCHIMIC	4.580	3.932				
9	TALANTA	4.162	3.841				
10	CRITICAL REV	4.000	4.301				
11	JOURNAL OF	3.981	4.008				
12	ANALYST	3.885	3.865				
13	Environmental	3.516	3.500				
14	JOURNAL OF	3.471	4.152				
15	Drug Testing a	3.469	2.694				
16	ANALYTICAL	3.431	3.306				
17	JOURNAL OF	3.379	3.205				
18	JOURNAL OF PHARMACEUTICAL AND BIOMEDICAL ANALYSIS	3.255	2.953				
19	MICROCHEMICAL JOURNAL	3.034	3.213				
20	JOURNAL OF ELECTROANALYTICAL CHEMISTRY	3.012	2.883				

数据视图 变量视图

添加变量(V)... IBM SPSS Statistics Processor 就绪

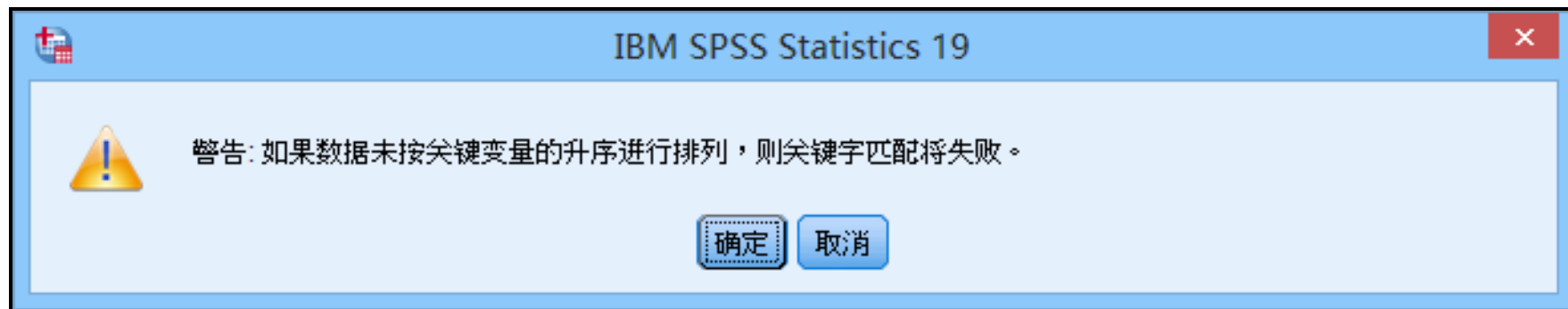
数据合并

□两个文件的数据变量中“期刊名称”是统一的，
因此可以将“期刊名称”作为关键变量



数据合并

□合并前还需要将关键变量进行升序排序



数据拆分

□其功能为按某一个变量值进行分组，数据仍在同一个文件中。但是，以后进行统计分析时，将根据拆分结果分别进行统计。

- 例5：对例5. sav文件数据：2010-2017清华大学、北京大学和四川大学在CNS上发表的论文信息（与实际数据略有删减）进行拆分
- 目的是分别查看三本期刊中三所高校在发表论文数量上的分布

*例5.sav [数据集2] - IBM SPSS

文件(F) 编辑(E) 视图(V) 数据(D) 转换(T) 分析(A) 直销(M) 图形(G) 实用程序(U) 窗口

4:

	论文标题	年份	机构
1	Whole-ge...	2017	北京大学
2	Ubiquitin...	2017	
3	Ubiquitin...	2017	
4	Transbou...	2017	
5	Transbou...	2017	
6	Tough ad...	2017	
7	The funda...	2017	
8	The close...	2017	
9	The Apos...	2017	
10	Synthesi...	2017	
11	Structure...	2017	
12	Structure...	2017	
13	Structure...	2017	
14	Structure...	2017	
15	Structure...	2017	
16	Structure...	2017	
17	Structure...	2017	
18	Structure...	2017	
19	Structure...	2017	
20	Structure... Shen, HZ; Zho...	2017	SCIENCE
21	Structure... Li, MH; Zhou, ...	2017	NATURE
22	Spin-orbit... Kolkowitz, S; ...	2017	NATURE
23	Sleight of... Casasanto, D	2017	SCIENCE

定义变量属性(V)...
设置未知测量级别(L)...
复制数据属性(C)...
新建设定属性(B)...
定义日期(E)...
定义多重响应集(M)...
验证(L)
标识重复个案(U)...
标识异常个案(I)...
排序个案...
排列变量...
转置(N)...
合并文件(G)
重组(R)...
分类汇总(A)...
正交设计(H)
复制数据集(D)
拆分文件(F)...
选择个案...
加权个案(W)...

分割文件

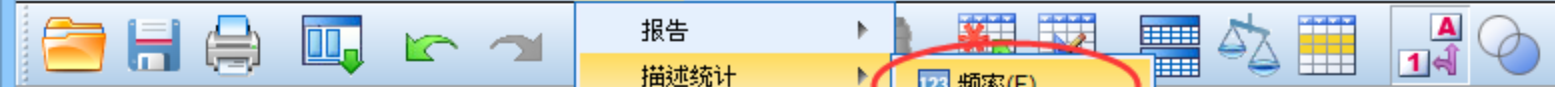
分组变量将安置在同一个表格中比较输出

分组变量将安置在不同的表格中比较输出

分析所有个案，不创建组(A)
比较组(C)
按组组织输出(O)
按分组变量排序文件(S)
文件已排序(F)

当前状态：按组合分析关闭。

确定 粘贴(P) 重置(R) 取消 帮助



7:

	论文标题	作者	
1	Structure...	Bai, R; Yan, ...	CEL
2	Structure...	Yan, Z; Zhou, ...	CEL
3	Structure...	Qian, HW; Zh...	CEL
4	Structure...	Wan, RX; Yan...	CEL
5	Structure...	Li, NN; Wu, J...	CEL
6	Structure...	Li, NN; Wu, J...	CEL
7	Regulator...	Wang, S; Xia,...	CEL
8	Nuclear ...	Lim, J; Giri, P...	CEL
9	Modeling ...	Chen, YC; Ch...	CEL
10	Landscap...	Zheng, CH; Zh...	CEL
11	In Situ C...	Liu, X; Zhang, ...	CEL
12	Derivation...	Yang, Y; Yan...	CEL
13	Derivation...	Yang, Y; Yan...	CEL
14	Architect...	Guo, RY; Zon...	CEL
15	An Atomi...	Zhang, XF; Ya...	CEL
16	3D Chro...	Ke, YW; Ke, ...	CEL
17	TMCO1 I...	Wang, QC; W...	CEL
18	Structure...	Wu, M; Gu, J...	CEL
19	Structural...	Gong, X; Qian...	CEL
20	RNA Dup...	Lu, ZP; Zhang...	CEL
21	Presynap...	Zhang, J; Zha...	CELL
22	Presvnap...	Zhano, J; Zha...	CELL

频率(F)

可见：5 变量的 5

变量	变量	变

论文标题

作者

期刊

年份

机构

☒ 显示频率表格(D)

统计量(S)...
图表(C)...
格式(F)...
Bootstrap(B)...

确定 粘贴(P) 重置(R) 取消 帮助

清华大学			
清华大学			
清华大学	2016		
北京大学	2016		

□输出结果

机构

期刊			频率	百分比	有效百分比	累积百分比
CELL	有效	北京大学	28	41.2	41.2	41.2
		清华大学	35	51.5	51.5	92.6
		四川大学	5	7.4	7.4	100.0
		合计	68	100.0	100.0	
NATURE	有效	北京大学	77	37.6	37.6	37.6
		清华大学	112	54.6	54.6	92.2
		四川大学	16	7.8	7.8	100.0
		合计	205	100.0	100.0	
SCIENCE	有效	北京大学	71	51.1	51.1	51.1
		清华大学	63	45.3	45.3	96.4
		四川大学	5	3.6	3.6	100.0
		合计	139	100.0	100.0	

数据排序

□将数据按指定的某一个或多个变量值的升序或降序重新排列，所指定的变量称为排序变量。

- 单值排序：排序变量只有一个。
- 多重排序：排序变量有多个，多重排序的第一个排序变量称为主排序变量，其它排序变量依次称为第二排序变量、第三排序变量等。
- 例6：打开例6. sav数据文件。要求职称按升序排序，工资按降序排序。

例6.sav [数据集1] - IBM SPSS Statistics 数据编辑器

文件(F) 编辑(E) 视图(V) 数据(D) 转换(T) 分析(A) 直销(M) 图形(G) 实用程序(U) 窗口(W) 帮助

6:

	职工号	性别	变量
1	4	男	
2	1	男	
3	2	女	
4	3	男	
5	6	女	
6	7	女	
7	8	男	
8	9	女	
9	5		
10			
11			
12			
13			
14			
15			
16			
17			
18			
19			
20			
21			
22			
23			
24			

定义变量属性(V)...
设置未知测量级别(L)...
复制数据属性(C)...
新建设定属性(B)...
定义日期(E)...
定义多重响应集(M)...
验证(L)
标识重复个案(U)...
标识异常个案(I)...
排序个案...
排列变量...
转置(N)...
合并文件(G)
重组(R)...
分类汇总(A)...
正交设计(H)
复制数据集(D)
拆分文件(F)...
选择个案...
加权个案(W)...

排序个案...

数据视图 变量视图

排序个案...

IBM SPSS Statistics Processor 就绪

排序个案

排序依据:

- 职工号
- 性别
- 年龄
- 职称(A)
- 工资(D)

排列顺序:

☐ 升序(A)

☒ 降序(D)

确定 粘贴(P) 重置(R) 取消 帮助

数据的分类汇总

□按指定的分类变量对观测值进行分组，对每组记录的各变量求指定的描述统计量。

- 例7：将例7. sav中数据，以职称作为分组变量，对职工年龄的均值和工资的中值进行汇总。

例7.sav

文件(F) 编辑(E) 视图(V) 数据(D) 转换(T) 分析(A) 直销(M) 图

2:

	职工号	性别
1	1	0
2	2	0
3	3	0
4	4	0
5	5	0
6	6	0
7	7	0
8	8	0
9	9	0
10	10	0
11	11	0
12	12	0
13	13	0
14	14	0
15	15	0
16	16	0
17	17	0
18	18	0
19	19	0
20	20	0
21	21	0
22	22	0

定义变量属性(V)...
设置未知测量级别(L)...
复制数据属性(C)...
新建设定属性(B)...
定义日期(E)...
定义多重响应集(M)...
验证(L)
标识重复个案(U)...
标识异常个案(I)...
排序个案...
排列变量...
转置(N)...
合并文件(G)
重组(R)...
分类汇总(A)...
正交设计(H)
复制数据集(D)
拆分文件(F)...
选择个案...
加权个案(W)...

汇总数据

分组变量(B): 职称

汇总变量

变量摘要(S):
年龄_mean = MEAN(年龄)
工资_median = MEDIAN(工资)

函数(F)... 变量名与标签(N)...
☐ 个案数(C) 名称(M): N_BREAK

保存

☒ 将汇总变量添加到活动数据集(D)
☐ 创建只包含汇总变量的新数据集(E)
数据集名称(D):
☒ 写入只包含汇总变量的新数据文件(W)
文件(L)... C:\Users\501\Documents\laggr.sav

适用于大型数据集的选项

☐ 文件已经按分组变量排序(A)
☐ 在汇总之前排序文件(G)

确定 粘贴(P) 重置(R) 取消 帮助

数据的加权

□ 加权是为了告知统计软件你这一行数据代表的并不是单个值，而是表示实际样本很多个，有相应的“频数”之和那么多的样本数。

- 例8：对例8.sav数据进行加权，分析我校在2017年化学学科发表的期刊中，1-4区期刊的占比。

*未标题2 [数据集1] - IBM SPSS Statistics

文件(F) 编辑(E) 视图(V) 数据(D) 转换(T) 分析(A) 直销(M) 图形(G) 实用程序(U) 窗口(W)

1: 分区 Q2

	论文数量	分区
1	120	Q2
2	36	Q1
3	35	Q2
4	27	Q1
5	27	Q1
6	26	Q1

定义变量属性(V)...
设置未知测量级别(L)...
复制数据属性(C)...
新建设定属性(B)...
定义日期(E)...
定义多重响应集(M)...
验证(L)
标识重复个案(U)...
标识异常个案(I)...
排序个案...
排列变量...
转置(N)...
合并文件(G)
重组(R)...
分类汇总(A)...
正交设计(H)
复制数据集(D)
拆分文件(F)...
选择个案...
加权个案(W)... (highlighted)

加权个案

☐ 请勿对个案加权(D)
☒ 加权个案(W)

频率变量(F):
论文数量

当前状态: 不加权个案

确定 粘贴(P) 重置(R) 取消 帮助

IBM SPSS Statistics 26.0.0.0

*未标题2 [数据集1] - IBM SPSS Stat

文件(F) 编辑(E) 视图(V) 数据(D) 转换(T) 分析(A) 直销(M) 图形(G) 实用程序(U) 窗口(W)

1: 分区 Q2

	期刊名称
1	RSC ADVANCES
2	CHEMICAL COMMUNICATIONS
3	JOURNAL OF APPLIED POLYMER
4	ORGANIC LETTERS
5	POLYMER
6	INDUSTRIAL & ENGINEERING C
7	MACROMOLECULAR RAPID CO
8	ACS SUSTAINABLE CHEMISTRY
9	PHYSICAL CHEMISTRY CHEMIC
10	CHINESE CHEMICAL LETTERS
11	ANGEWANDTE CHEMIE-INTERN
12	CARBOHYDRATE POLYMERS
13	CHEMISTRY-A EUROPEAN JOU
14	JOURNAL OF COLLOID AND INT
15	POLYMER CHEMISTRY
16	INORGANIC CHEMISTRY
17	MICROCHEMICAL JOURNAL
18	JOURNAL OF POLYMER RESEA

报告

描述统计

表(T)

比较均值(M)

一般线性模型(G)

广义线性模型

混合模型(X)

相关(C)

回归(E)

对数线

神经网络

分类(F)

降维

度量(S)

非参数

预测(T)

生存函

多重响

缺失值分析(Y)...

多重归因(T)

复杂抽样(L)

频率(F)...

描述(D)...

探索(E)

交叉表(C)

比率(R)...

P-P 图...

频率

变量(V):

期刊名称

论文数量

VAR00001

VAR00002

VAR00003

VAR00004

VAR00005

VAR00006

VAR00007

分区

统计量(S)...

图表(C)...

格式(F)...

Bootstrap(B)...

显示频率表格(D)

确定 粘贴(P) 重置(R) 取消 帮助

	13	Q1
多重归因(T)	13	Q1
复杂抽样(L)	13	Q2

□结果

分区

		频率	百分比	有效百分比	累积百分比
有效	#N/A	7	.6	.6	.6
	Q1	607	49.3	49.3	49.9
	Q2	398	32.3	32.3	82.2
	Q3	135	11.0	11.0	93.2
	Q4	84	6.8	6.8	100.0
	合计	1231	100.0	100.0	

重复数据的查找

□定位重复的个体。适用于数据双录入后的数据检索。

- 例9：查找例9. sav中的重复数据。
- 数据中第1-500条个案是从WOS中下载的500篇论文，包括标题、WOS号、DOI号等信息。第501-805条个案是所有工程类论文的列表。
- 目的是找出第1-500条个案中属于工程类的论文。

文件(F) 编辑(E) 视图(V) 数据(D) 转换(T) 分析(A) 直销(M) 图形(G) 实用程序(U) 窗口

定义变量属性(V)...

设置未知测量级别(L)...

复制数据属性(C)...

新建设定属性(B)...

定义日期(E)...

定义多重响应集(M)...

验证(L)

标识重复个案(U)...

标识异常个案(I)...

排序个案...

排列变量...

转置(N)...

合并文件(G)

重组(R)...

分类汇总(A)...

正交设计(H)

复制数据集(D)

拆分文件(F)...

选择个案...

加权个案(W)...

1	Effective coating	mina via atomic layer deposi...	W
2	Citric acid indu	the support and enhanced c...	W
3	Facile syntheses	as advanced electrode mat...	W
4	Chromium (VI)	g porous magnetite nanoparti...	W
5	Freeze-drying	g P and Si and its flame reta...	W
6	Preparation of	uff powder concrete at the Da...	W
7	Study of the k	remagnetization: Implication...	W
8	Strengthening	ion	W
9	Invariant tori fo	ctive liquid crystals bearing p...	W
10	Synthesis and	sites with segregated nickel ...	W
11	Ultrahigh mole	bioink with desirable moldabi...	W
12	Degradation re	paper with high thermal sta...	W
13	Acid-base syn	ceptors on cellulose/1-allyl-...	W
14	An insight into	dation operated close to stoic...	W
15	Pd or PdO: Ca	phene surface: An all-electron...	W
16	Adsorption of	brid superparamagnetic nano...	W
17	Polymer-entan	tional graphene oxide as an e...	W
18	A Schiff base/	microspheres and their cataly...	W
19	Preparation of		
20	Synergistic effect of expandable graphite and melamine phosphate on flame...		W
21	Structure-Performance Relationships of Hole-Transporting Materials in Perov...		W
22	The Euclidean embedding learning based on convolutional neural network for ...		W
23	MMP-2 and Notch signal pathway regulate migration of adipose-derived stem...		W
24	Mixed mode fracture analysis of CCBD specimens based on the extended m...		W
25	Investigation of phase structure, microstructure, and electrical properties of L...		W
26	A TDDFT Investigation on Plasmons in Multilayer Graphene Nanostructures		W
27	Nanosized CeO2 Particles Obtained by Mechanical Solid-State Reaction Co...		W

数据视图

变量视图

IBM SPSS Statistics 数据编辑器

标识重复的个案

定义匹配个案的依据(D):
WOS号
DOI号
期刊
出版年

在匹配组内的排序标准(O):

排序
☒ 升序(C)
☐ 降序(E)

匹配和分类变量数: 1

要创建的变量
☒ 基本个案指示符 (1=唯一或基本, 0=重复) (I)
☐ 每组中的最后一个个案为基本个案(L)
☒ 每组中的第一个个案为基本个案(H)
☐ 根据指示符的值进行筛选(F)
☒ 连续计算每个组合中的匹配个案 (0=非匹配个案)
☒ 将匹配个案移至文件顶端(A)
☒ 显示已创建变量的显示频率(V)

名称(N): 工程类论文
名称(M): 匹配顺序

确定 粘贴(P) 重置(R) 取消 帮助

标识重复个案(U)...

IBM SPSS Statistics Processor 就绪

个案的选择

□根据不同的要求，从所有个案中筛选出特定的个案。可以通过给数据表设置选择条件或者过滤条件来满足这一要求。

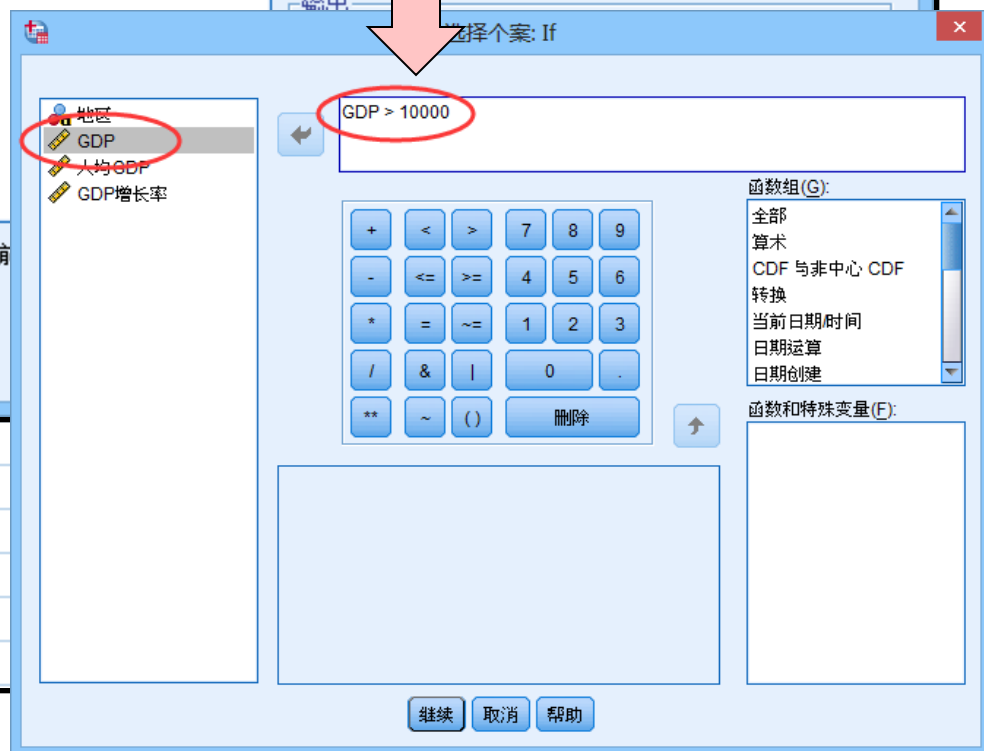
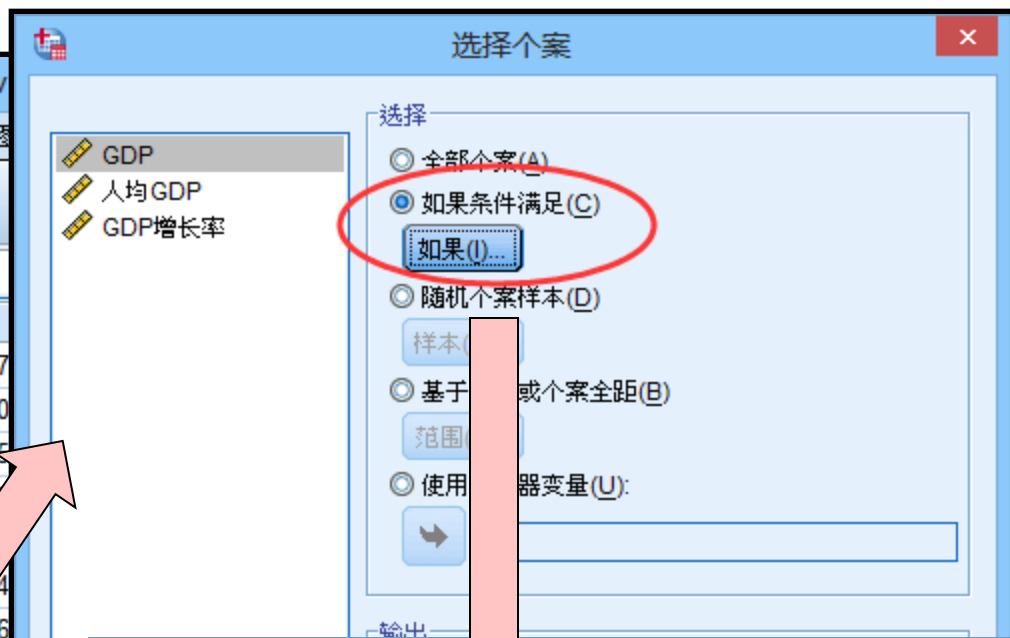
- 按条件选择：给出一个条件表达式，选取符合该表达式的个案
- 按数据范围选择：选择一定的数据范围内的全部个案，要求给出数据范围的上、下界个案编号
- 随机选择：对数据编辑窗口中的所有个案进行随机筛选
- 过滤变量选择：选择制定的一个已存在的变量作为个案选取的标准

□例10-1：选择GDP大于10000亿元的地区

□例10-2：选择GDP增长率在6%-9%之间的地区

例10.sav

	地区		DP
1	北京		106497
2	天津		107960
3	河北		40255
4	山西		3491
5	内蒙古		71
6	辽宁		4
7	吉林		086
8	黑龙江		39462
9	上海		103796
10	江苏		87995
11	浙江		77644
12	安徽		35997
13	福建		67966
14	江西		36724
15	山东		64168
16	河南		39123
17	湖北		50654
18	湖南		42754
19	广东		67503
20	广西	16803.12	35190
21	海南	3702.76	40818
22	重庆	15717.27	52321



计算新变量

□使用SPSS算数表达式及函数，对所有记录或满足SPSS条件表达式的记录，计算出一个新结果，并将结果存入一个指定的变量中。

- 例11：计算例11. sav文件数据中男生的平均成绩。

例11.sav [数据集1] - IBM SPSS Statistics 数据编辑器

文件(F) 编辑(E) 视图(V) 数据(D) 转换(T) 分析(A) 直销(M) 图形(G) 实用程序(U) 窗口(W) 帮助



计算变量(C)...

对个案内的值计数



1:

	编号	性别
1	1	
2	2	
3	3	
4	4	
5	5	
6	6	
7	7	
8	8	
9	9	
10	10	
11		
12		
13		
14		
15		
16		
17		

目标变量(T):

平均成绩

数字表达式(E):

MEAN(数学,英语,语文)

类型与标签(L)...

- 编号
- 性别
- 数学
- 英语
- 语文

如果(I)... (可选的个)

计算变量

计算变量: If 个案

☐ 包括所有个案(A)

☒ 如果个案满足条件则包括(F):

性别 = 1



函数组(G):

全部
算术
CDF 与非中心 CDF
转换
当前日期/时间
日期运算
日期创建

函数和特殊变量(F):

\$Casenum
\$Date
\$Date11
\$JDate
\$Sysmis
\$Time
Abs
Any
Applymodel
Arsin

继续 取消 帮助

数据视图 变量视图

计算变量(C)...

变量值的重新编码

□数据分析中，将连续变量转换为登记变量，或者将分类变量不同的变量等级进行合并，例如把同学的成绩分为优、良、中、差四个等级。

- 重新编码为相同变量：对原始变量的取值进行修改，用新编码直接取代原变量的取值。
- 重新编码为不同变量：将新编码存入新的变量，根据原始变量的取值生成一个新变量来表示分组情况。

□例12：将例12.sav中论文按被引次数分成三个等级。

- 论文发表年份：2015；所在学科：材料科学；

- 引用次数基准线：

- $15 < \text{citation}$ ：前10%，高水平论文

- $9 < \text{citation} \leq 15$ ：10%-50%，优秀论文

- $\text{Citation} \leq 9$ ：50%-，一般论文

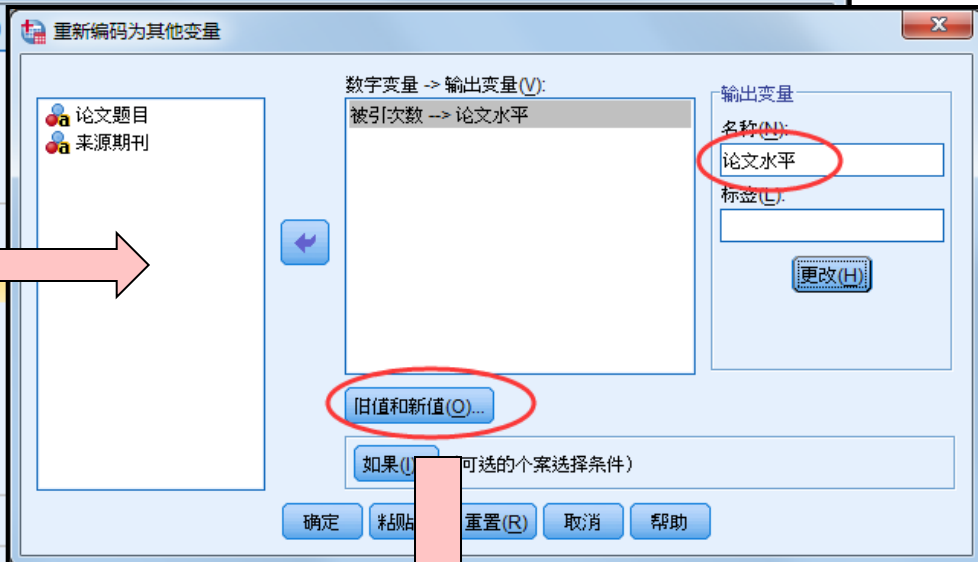
文件(F) 编辑(E) 视图(V) 数据(D) 转换(T) 分析(A) 直销(M) 图形(G)



1:

	论
1	antum oscillations in a t
2	existence of supercondu
3	ge-area high-quality 2D u
4	cile Synthesis of Ultrasm
5	02 nanotube arrays base
6	earable electrode-free trib
7	rous Two-Dimensional Na
8	tional Design of Small M
9	minescence-Driven Rever
10	Gylated Polypyrrole Nan
11	cent progress in piezoele
12	rsatile third components
13	electrocatalytic Hydrogen E
14	signing Efficient Non-Fullerene Acceptors by Tailoring ...
15	roporous free-standing nano-sulfur/reduced graphene ...
16	cent advances in antireflective surfaces based on nano...
17	porous nitrogen and phosphorous dual doped graphene ...
18	ll-to-Roll Green Transfer of CVD Graphene onto Plastic ...
19	ree-Dimensional Nitrogen-Doped Graphene Nanoribbon...
20	Iron-based Film for Highly Efficient Electrocatalytic Ox...
21	hene Quantum Dot Doping of MoS2 Monolayers

- 计算变量(C)...
- 对个案内的值计数(O)...
- 转换值(F)...
- 重新编码为相同变量(S)...
- 重新编码为不同变量(R)...
- 自动重新编码(A)...
- 可视离散化(B)...
- 最优离散化(I)...
- 准备建模数据(P)
- 个案排序(K)...
- 日期和时间向导...
- 创建时间序列(M)...
- 替换缺失值(V)...
- 随机数字生成器(G)...
- 运行挂起的转换(T)



数据视图 变量视图

重新编码为不同变量(R)...

主要内容

□1-SPSS概述

□2-SPSS数据管理

□3-SPSS的统计分析功能

3-SPSS的统计分析功能

□3.1-T检验

□3.2-方差分析

□3.3-线性回归与相关

□3.4-聚类分析

□3.5-因子分析

3.1-T检验

□T检验是检验样本的均值和给定的均值是否存在显著性差异。T检验分为三类：

□单样本检验（单个总体，方差未知，均值的检验）

□独立样本检验（两个总体，方差未知但相等，均值是否相等的检验）

□配对样本检验

单样本检验

□总体分布为正态分布 $N(\mu, \sigma^2)$ 时，需要检验

$$H_0: \mu = \mu_0.$$

- 检验统计量 $t = \frac{\bar{x} - \mu}{s} \sqrt{n}$
- SPSS将自动把样本均值 μ_0 、样本方差、样本数带入上式，计算出t统计量的观测值和对应的概率P值。
- 如果概率P值小于显著性水平 α ，则拒绝原假设，认为总体均值与检验值之间存在显著差异；反之则接受原假设。

单样本检验

□例13：某药物在某种溶剂中溶解后的标准浓度为20.00mg/L，现采用某种方法，测量该药物溶解度11次，测量后得到的结果见例13.sav，问：用该方法测量所得结果是否与标准浓度值有所不同？

单样本检验

例13.sav [数据集1] - IBM SPSS Statistics 数据编辑器

文件(F) 编辑(E) 视图(V) 数据(D) 转换(T) 分析(A) 直销(M) 图形(G) 实用程序(U) 窗口(W) 帮助

报告
描述统计
表(T)
比较均值(M)
一般线性模型(G)
广义线性模型
混合模型(X)
相关(C)
回归(R)

均值(M)...
单样本 T 检验(S)...
独立样本 T 检验(T)...
配对样本 T 检验(P)...
单因素 ANOVA...

可见：1 变量的 1

	浓度	变量	变量
1	20.99		
2	20.41		
3	20.10		
4	20.00		
5	20.91		
6	22.41		

单样本 T 检验

检验变量(T): 浓度

检验值(V): 20.00

选项(O)...
Bootstrap(B)...

单样本 T 检验: 选项

置信区间百分比(C): 95 %

缺失值
☒ 按分析顺序排除个案(A)
☐ 按列表排除个案(L)

继续 取消 帮助

数据视图 变量视图

质量控制(Q)
☒ ROC 曲线图(V)...

单样本 T 检验(S)... IBM SPSS Statistics Processor 就绪

单样本检验

结果显示

- $P\text{值}=0.012 < 0.05$ ，因此认为测量所得结果与标准浓度值有差异

→ T检验

[数据集1] I:\数据\例13.sav

单个样本统计量

	N	均值	标准差	均值的标准误
浓度	11	20.9836	1.06750	.32186

单个样本检验

	检验值 = 20.00					
	t	df	Sig.(双侧)	均值差值	差分的 95% 置信区间	
					下限	上限
浓度	3.056	10	.012	.98364	.2665	1.7008

独立样本检验

□两个独立样本符合正态分布，且满足方差齐性。

需要检验 $H_0: \mu_1 - \mu_2 = 0$ 。

- 选取检验统计量为t统计量，
$$t = \frac{\bar{X}_1 - \bar{X}_2 - (\mu_1 - \mu_2)}{\sqrt{\frac{S_1^2}{n_1} + \frac{S_2^2}{n_2}}}$$
- 计算F统计量和t统计量的观测值以及相应的概率P值。
- 利用F检验判断两总体的方差是否相等。
- 利用t检验判断两总体的均值是否存在显著差异。

独立样本检验

□例14：现希望评价两位老师的教学质量，试比较其分别任教的两班考试后的成绩是否存在差异。数据见例14.sav。

独立样本检验

例14.sav [数据集1] - IBM SPSS Statistics 数据编辑器

文件(F) 编辑(E) 视图(V) 数据(D) 转换(T) 分析(A) 直销(M) 图形(G) 实用程序(U) 窗口(W) 帮助

报告
描述统计
表(T)
比较均值(M)
一般线性模型(G)
广义线性模型
混合模型(X)
相关(C)
回归(R)
对数线性模型(O)
神经网络
分类(F)
降维
度量(S)
非参数检验(N)
预测(T)
生存函数(S)
多重响应(U)
缺失值分析(Y)...
多重归因(T)
复杂抽样(L)
质量控制(Q)
ROC 曲线图(V)...

均值(M)...
单样本 T 检验(S)...
独立样本 T 检验(T)...
配对样本 T 检验(P)...
单因素 ANOVA...

可见：2 变量的 2

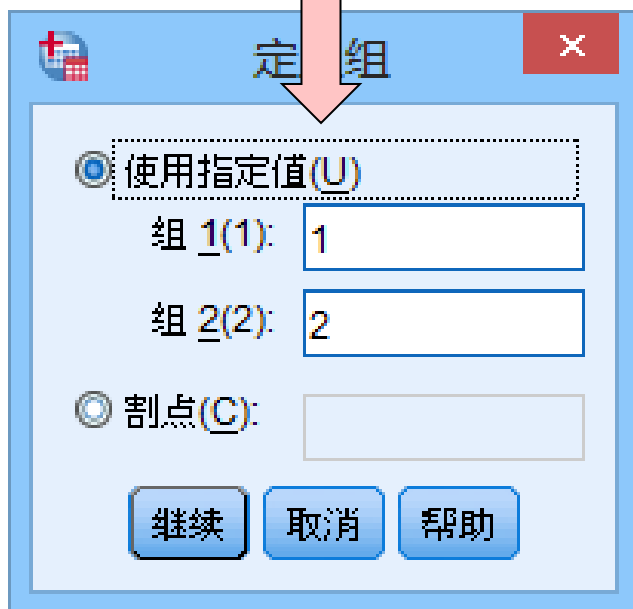
	score	class	变
1	85	甲班	
2	73	甲班	
3	86	甲班	
4	77	甲班	
5	94	甲班	
6	68	甲班	
7	82	甲班	
8	83	甲班	
9	90	甲班	
10	88	甲班	
11	76	甲班	
12	85	甲班	
13	87	甲班	
14	74	甲班	
15	85	甲班	
16	80	甲班	
17	82	甲班	

数据视图 变量视图

独立样本 T 检验(T)...

IBM SPSS Statistics Processor 就绪

独立样本检验



“使用指定值”表示用实现定义好的变量值表示不同的分组，在本例中，甲班的值为1，乙班的值为2；

“割点”表示分组变量为连续变量时，输入一个数字，大于等于该数值的为一个总体，对应一组样本，小于该值的为另一总体，对应另一组样本

独立样本检验

结果显示

- F检验的P值为0.397>0.05，故方差齐性。
- 不同组间独立样本t检验P值为0.004<0.05，因此认为甲、乙两班的成绩存在差异。

→ T检验

[数据集1] I:\数据\例14.sav

组统计量

	class	N	均值	标准差	均值的标准误
score	甲班	20	83.30	6.906	1.544
	乙班	20	75.45	9.179	2.053

独立样本检验

		方差方程的 Levene 检验		均值方程的 t 检验						
		F	Sig.	t	df	Sig.(双侧)	均值差值	标准误差值	差分的 95% 置信区间	
									下限	上限
score	假设方差相等	.733	.397	3.056	38	.004	7.850	2.569	2.650	13.050
	假设方差不相等			3.056	35.290	.004	7.850	2.569	2.637	13.063

配对样本检验

- 利用来自两个不同总体的配对样本，推断两个总体的均值是否有差异。
- 对两组样本分别计算每对观察值的差值得到差值样本，然后利用差值样本，通过对其均值是否显著为0的检验来推断两总体均值的差是否显著为0。
- 例15：某地区随机抽取12名贫血儿童的家庭，实行健康教育干预三个月，干预前后儿童的血红蛋白（%）测量结果见例15.sav，试问干预前后该地区贫血儿童血红蛋白（%）平均水平有无变化？

配对样本检验

例15.sav [数据集2] - IBM SPSS Statistics 数据编辑器

文件(F) 编辑(E) 视图(V) 数据(D) 转换(T) 分析(A) 直销(M) 图形(G) 实用程序(U) 窗口(W) 帮助

报告
描述统计
表(T)
比较均值(M)
一般线性模型(G)
广义线性模型
混合模型(X)
相关(C)
回归(R)
对数线性模型(O)

均值(M)...
单样本 T 检验(S)...
独立样本 T 检验(T)...
配对样本 T 检验(P)...
单因素 ANOVA...

可见：3 变量的 3

	序号	干预前	干预后
1	1	36	
2	2	46	
3	3	53	
4	4	57	
5	5	65	
6	6	60	
7	7	42	
8	8	45	
9	9	25	
10	10	55	
11	11	51	
12	12	59	
13			
14			
15			
16			
17			

数据视图 变量视图

配对样本 T 检验(P)...

配对样本 T 检验

成对变量(V):

对(A)	Variable1	Variable2
1	[干预前]	[干预后]
2		

确定 粘贴(P) 重置(R) 取消 帮助

配对样本检验

结果显示

- 统计量P值为 $0.007 < 0.05$ ，因此认为干预前后该地区贫血儿童血红蛋白（%）水平有变化。

成对样本检验									
		成对差分				t	df	Sig.(双侧)	
		均值	标准差	均值的标准误	差分的 95% 置信区间				
					下限				上限
对 1	干预前 - 干预后	-10.667	11.179	3.227	-17.769	-3.564	-3.305	11	.007

3.2-方差分析

- 当比较两组资料均值是否相等时，可以采用t检验。
当组数大于2组时，如果仍然采用t检验，这不但复杂，而且有很大的可能性导致错误结论。这时应该采用方差分析。
- 方差分析的应用条件如下：独立；正态；方差齐性。

□例16：比较3个不同电池生产企业生产电池的寿命，见例16.sav

例16.sav [数据集1] - IBM SPSS Statistics 数据编辑器

文件(F) 编辑(E) 视图(V) 数据(D) 转换(T) 分析(A) 直销(M) 图形(G) 实用程序(U) 窗口(W) 帮助



	企业	寿命	变量
1	1	40	
2	1	48	
3	1	38	
4	1	42	
5	1	45	
6	1	43	
7	1	42	
8	1	39	
9	1	48	
10	1	44	
11	1	47	
12	1	43	
13	2	26	
14	2	31	
15	2	30	
16	2	34	
17	2	34	

数据视图 变量视图

单因素 ANOVA...

报告

描述统计

表(T)

比较均值(M)

一般线性模型(G)

广义线性模型

混合模型(X)

相关(C)

回归(R)

对数线性模型(O)

神经网络

分类(F)

降维

度量(S)

非参数检验(N)

预测(T)

生存函数(S)

多重响应(U)

缺失值分析(Y)

多重归因(T)

复杂抽样(L)

质量控制(Q)

ROC 曲线图(V)

均值(M)...

单样本 T 检验(S)...

独立样本 T 检验(T)...

配对样本 T 检验(P)...

单因素 ANOVA...

可见：2 变量的 2

变量

单因素方差分析

因变量列表(E):

电池 [寿命]

对比(N)...

两两比较(H)...

选项(O)...

Bootstrap(B)...

因子(F):

企业

确定

粘贴(P)

重置(R)

取消

帮助

IBM SPSS Statistics Processor 就绪

□两两比较：如果结果显示不同企业生产的电池寿命存在差异，那么通过“两两比较”可以获知是哪两个厂家的电池有差异。

单因素 ANOVA: 两两比较

假定方差齐性

☒ LSD(L) ☒ S-N-K(S) ☐ Waller-Duncan(W)

☐ Bonferroni(B) ☐ Tukey 类型 I/类型 II 误差比率(I): 100

☐ Sidak ☐ Tukey s-b(K) ☐ Dunnett(E)

☐ Scheffe(C) ☐ Duncan(D) 控制类别: 最后一个(L)

☐ R-E-G-W F(R) ☐ Hochberg's GT2(H) 检验

☐ R-E-G-W Q(Q) ☐ Gabriel(G) ☒ 双侧(2) ☐ < 控制(O) ☐ > 控制(N)

未假定方差齐性

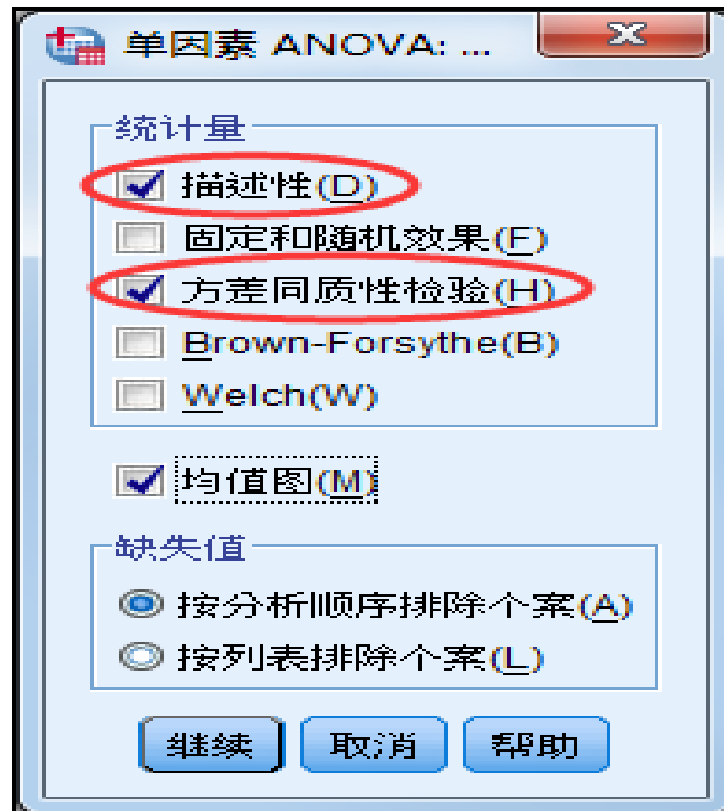
☐ Tamhane's T2(M) ☐ Dunnett's T3(3) ☐ Games-Howell(A) ☐ Dunnett's C(U)

显著性水平(F): 0.05

继续 取消 帮助

□选项

- 描述性：显示每个因变量的个数、均值、标准差等信息
- 方差同质性检验：计算Levene方差齐性检验



结果显示-1

描述

电池

	N	均值	标准差	标准误	均值的 95% 置信区间		极小值	极大值
					下限	上限		
1	12	43.25	3.334	.962	41.13	45.37	38	48
2	12	32.33	3.576	1.032	30.06	34.61	26	37
3	12	43.83	3.881	1.120	41.37	46.30	39	50
总数	36	39.81	6.405	1.067	37.64	41.97	26	50

□结果显示-2

- 方差齐性检验：显著性=0.680>0.05，各组方差齐性

方差齐性检验			
电池			
Levene 统计量	df1	df2	显著性
.390	2	33	.680

结果显示-3

- 显著性=0.000<0.05，表示三个厂家生产的电池寿命存在差异

ANOVA

电池

	平方和	df	均方	F	显著性
组间	1007.056	2	503.528	38.771	.000
组内	428.583	33	12.987		
总数	1435.639	35			

结果显示-4 (LSD法结果)

- 企业1与企业2显著性=0.000<0.05, 存在差异
- 企业1与企业3显著性=0.694>0.05, 无差异
- 企业2与企业3显著性=0.000<0.05, 存在差异

多重比较

因变量: 电池

		均值差 (I-J)	标准误	显著性	95% 置信区间	
(I) 企业	(J) 企业				下限	上限
LSD 1	2	10.917 [*]	1.471	.000	7.92	13.91
	3	-.583	1.471	.694	-3.58	2.41
2	1	-10.917 [*]	1.471	.000	-13.91	-7.92
	3	-11.500 [*]	1.471	.000	-14.49	-8.51
3	1	.583	1.471	.694	-2.41	3.58
	2	11.500 [*]	1.471	.000	8.51	14.49

*. 均值差的显著性水平为 0.05。

结果显示-5 (SNK法结果)

- 企业2是一个类别，企业1与企业3是一个类别

同类子集

电池

		N	alpha = 0.05 的子集	
			1	2
Student-Newman-Keuls ^a	企业 2	12	32.33	
	1	12		43.25
	3	12		43.83
	显著性		1.000	.694

将显示同类子集中的组均值。

a. 将使用调和均值样本大小 = 12.000。

3.3-线性回归与相关

□线性相关系数：Pearson相关系数

□取值范围 $-1 \leq r \leq 1$

□绝对值越接近1，表示两变量间的相关关系的密切程度越高

□例17：分析发文量、被引次数、h指数、篇均被引次数4个指标之间的相关性。相关数据见例17.sav。

3.3-线性回归与相关

*未标题2 [数据集1] - IBM SPSS Statistics 数据编辑器

文件(F) 编辑(E) 视图(V) 数据(D) 转换(T) **分析(A)** 直销(M) 图形(G) 实用程序(U) 窗口(W) 帮助

报告
描述统计
表(T)
比较均值(M)
一般线性模型(G)
广义线性模型
混合模型(X)
相关(C)
回归(R)
对数线性模型(O)
神经网络
分类(F)
降维
度量(S)
非参数检验(N)
预测(T)
生存函数(S)
多重响应(U)
缺失值分析(Y)...
多重归因(T)
复杂抽样(L)
质量控制(Q)
ROC 曲线图(V)...

可见: 5 变量的 5

	引次数	h指数	篇均被引
1	19610	59	
2	26614	67	
3	30572	75	
4		79	
5		65	
6		47	
7		36	
8	10353		
9	13899		
10	7532		
11	6650		
12	5222		
13	9273		
14	4280		
15	2385		
16	3336		
17	5761		

双变量(B)...
偏相关(R)...
距离(D)...

双变量相关

变量(V):
发文章
被引次数
h指数
篇均被引次数

相关系数
☒ Pearson ☐ Kendall's tau-b(K) ☐ Spearman

显著性检验
☒ 双侧检验(T) ☐ 单侧检验(L)

☒ 标记显著性相关(F)

确定 粘贴(P) 重置(R) 取消 帮助

数据视图 变量视图

双变量(B)...

IBM SPSS Statistics Process

3.3-线性回归与相关

结果显示

相关性

[数据集1]

相关性

		发文量	被引次数	h指数	篇均被引次数
发文量	Pearson 相关性	1	.947**	.913**	-.104
	显著性 (双侧)		.000	.000	.530
	N	39	39	39	39
被引次数	Pearson 相关性	.947**	1	.926**	-.032
	显著性 (双侧)	.000		.000	.847
	N	39	39	39	39
h指数	Pearson 相关性	.913**	.926**	1	-.147
	显著性 (双侧)	.000	.000		.372
	N	39	39	39	39
篇均被引次数	Pearson 相关性	-.104	-.032	-.147	1
	显著性 (双侧)	.530	.847	.372	
	N	39	39	39	39

** . 在 .01 水平 (双侧) 上显著相关。

3.3-线性回归与相关

□线性回归是分析两个定量变量间数量依存关系的统计分析方法。

□回归分析主要包括三方面内容

- 提供建立有相关关系的变量之间的数学关系式
- 判别影响变量的众多变量中哪些影响是显著的
- 利用所得到的经验公式进行预测和控制

□例18：对某省9个地区水质的碘含量及其甲状腺肿的患病率作调查得到一组数据，见例18.sav，试进行回归分析

3.3-线性回归与相关

例18.sav [数据集1] - IBM SPSS Statistics 数据编辑器

文件(F) 编辑(E) 视图(V) 数据(D) 转换(T) **分析(A)** 直销(M) 图形(G) 实用程序(U) 窗口(W) 帮助

报告
描述统计
表(T)
比较均值(M)
一般线性模型(G)
广义线性模型
混合模型(X)
相关(C)
回归(R)
对数线性模型
神经网络
分类(F)
降维
度量(S)
非参数检验(N)
预测(T)
生存函数(S)
多重响应(U)
缺失值分析(Y)
多重归因(T)
复杂抽样(L)
质量控制(Q)
ROC 曲线图(O)

	地区	碘含量	患
1	1	1.0	
2	2	2.0	
3	3	2.5	
4	4	3.5	
5	5	3.5	
6	6	4.0	
7	7	4.4	
8	8	4.5	
9	9	5.2	
10			
11			
12			
13			
14			
15			
16			
17			

数据视图 变量视图

线性(L)...

线性回归

因变量(D): 患病率 [患病率]

自变量(I): 碘含量 [碘含量]

方法(M): 进入

统计量(S)...
绘制(T)...
保存(S)...
选项(O)...
Bootstrap(B)...

选择变量(E):
个案标签(C):
WLS 权重(H):

确定 粘贴(P) 重置(R) 取消 帮助

3.3-线性回归与相关

模型汇总

模型	R	R 方	调整 R 方	标准 估计的误差
1	.971 ^a	.943	.934	1.5747

a. 预测变量: (常量), 碘含量。

1、线性回归出来的相关系数为 **$R=0.971$** 。

2、方程拟合优度 **R 方是0.943**，调整后的 **R 方为0.934**

R 方是对回归方程拟合情况的描述， **R 方**是方程中变量 **X** 对 **Y** 的解释程度，越接近**1**，表明方程中 **X** 对 **Y** 的解释能力越强，拟合度越好。

3.3-线性回归与相关

Anova^b

模型	平方和	df	均方	F	Sig.
1 回归	285.504	1	285.504	115.136	.000 ^a
残差	17.358	7	2.480		
总计	302.862	8			

a. 预测变量: (常量), 碘含量。

b. 因变量: 患病率

在确认线性回归之前，必须判断变量的关系是否满足一元线性模型，即转换由 $Y=a+bX=e$ ， e 服从正态分布，检验假设 $H_0: b=0$ ； $H_1: b \neq 0$ 。

F统计量P值=0<0.05，说明模型整体是显著的，具有统计学意义

3.3-线性回归与相关

系数^a

模型	非标准化系数		标准系数	t	Sig.
	B	标准 误差	试用版		
1 (常量)	17.484	1.507		11.600	.000
碘含量	4.459	.416	.971	10.730	.000

a. 因变量: 患病率

给出了常数项和回归系数。
患病率=17.484+4.459*碘含量

3.4-聚类分析

聚类是根据某些数量特征将观察对象进行分类的一种数理统计方法。

- **系统聚类法：**首先将 n 个样品看成 n 类，然后将性质最接近的两类合并为一类，得到 $n-1$ 类，然后再从这些类中找出性质最接近的两个类合并为 $n-2$ 类，重复上述步骤，一直到所有样品聚为一类。整个过程可以绘成聚类图或树状图。
- **动态分类法：**首先将样品粗糙分为 n 类，然后根据某种最优准则进行调整至不能调整为止。
-

3.4-聚类分析

- **K-中心聚类**：快速高效，特别是大量数据时，准确性高一些，但是需要指定聚类的类别数量。
- **例20**：根据30所大学的各指标数据（例20.sav），将其分为4类。

*未标题2 [数据集1] - IBM SPSS Statistics 数据编辑器

文件(F) 编辑(E) 视图(V) 数据(D) 转换(T) **分析(A)** 直销(M) 图形(G) 实用程序(U) 窗口(W) 帮助

报告
描述统计
表(T)
比较均值(M)
一般线性模型(G)
广义线性模型
混合模型(X)
相关(C)
回归(R)
对数线性模型(O)
神经网络
分类(E)
降维
度量(S)
非参数检验(N)
预测(T)
生存函数(S)
多重响应(U)
缺失值分析(Y)...
多重归因(T)
复杂抽样(L)
质量控制(Q)
ROC 曲线图(V)...

可见：9 变量的 9

	机构	被引次数	篇均被引次数	学科规范化引文印象里
1	Zhejiang University	848961	11.24	1.12
2	Shanghai Jiao Tong University	793936	10.83	1.14
3	Peking University	924786	13.52	1.34

K 均值聚类分析

变量(V):

- 篇均被引次数
- 学科规范化引文印象里
- 被引率
- 高被引率
- 前10%论文率 [前10论文率]
- 国际合作论文率

迭代(I)...
保存(S)...
选项(O)...

个案标记依据(B):

方法

☒ 迭代与分类(T) ☐ 仅分类(Y)

聚类中心

☐ 读取初始聚类中心(E):

☒ 打开数据集(N) 文件(F)...

☒ 外部数据文件(X)

☐ 写入最终聚类中心(W):

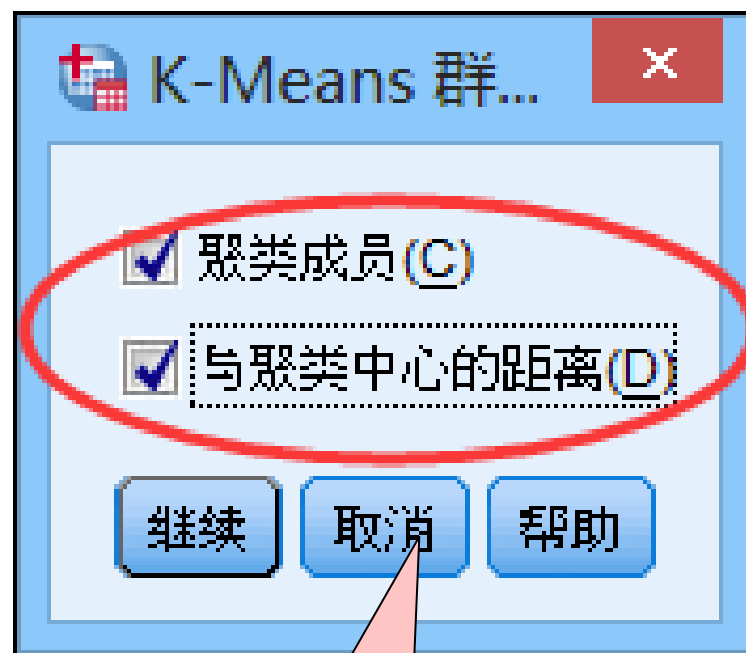
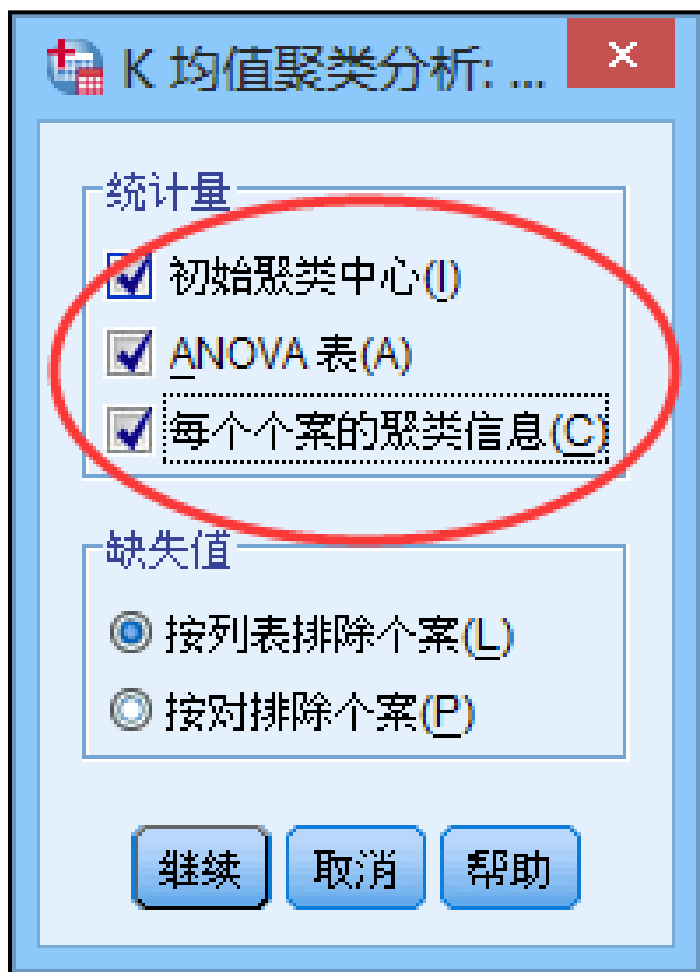
☒ 新数据集(D)

☒ 数据文件(A) 文件(L)...

确定 粘贴(P) 重置(R) 取消 帮助

数据视图 变量视图

K-均值聚类(K)...



在数据中新加两列表示样本所属的类别以及与聚类中心的距离

选取了其中4个样本作为初始聚类中心

初始聚类中心

	聚类			
	1	2	3	4
论文篇数	17248	21605	68393	37635
被引次数	127523	322748	924786	553011
篇均被引次数	7.39	14.94	13.52	14.69
学科规范化引文印象里	1.01	1.34	1.34	1.41
被引率	71.72	81.22	76.43	80.04
高被引率	.93	1.71	1.71	1.97
前10%论文率	9.49	14.84	13.13	15.33
国际合作论文率	24.15	22.47	34.89	30.32
	重庆大学	南开大学	北京大学	中科大

□最终聚类中心

最终聚类中心

	聚类			
	1	2	3	4
论文篇数	18874	31416	70506	44832
被引次数	166265	314972	860243	575029
篇均被引次数	8.80	10.31	12.27	12.92
学科规范化引文印象里	1.12	1.14	1.25	1.28
被引率	74.25	75.89	77.06	77.07
高被引率	1.15	1.15	1.46	1.43
前10%论文率	10.06	10.92	12.10	12.95
国际合作论文率	27.65	23.61	29.87	28.27

聚类结果

聚类成员		
案例号	聚类	距离
1	3	12345.772
2	3	66366.096
3	3	64577.828
4	3	14238.975
5	4	97716.315
6	4	47894.580
7	2	84661.599
8	2	96682.683
9	4	27674.896
10	2	100442.996
11	2	60920.591
12	4	23163.764
13	2	33152.465
14	2	4773.266
15	2	34247.604
16	2	23737.988
17	2	63562.157
18	2	22757.203
19	2	53529.304
20	2	72385.791
21	2	31806.703

□系统聚类：实际工作中使用较多的一种聚类方法，它可以对样品聚类，也可以对变量聚类。

□例21：根据例21.sav中数据，对指标进行分类。

***例20.sav [数据集1] - IBM SPSS Statistics 数据编辑器**

文件(F) 编辑(E) 视图(V) 数据(D) 转换(T) **分析(A)** 窗口(W) 帮助

报告
描述统计
表(T)
比较均值(M)
一般线性模型(G)
广义线性模型
混合模型(X)
相关(C)
回归(R)
对数线性模型(O)
神经网络
分类(F)
降维
度量(S)
非参数检验(N)
预测(T)
生存函数(S)
多重响应(U)
缺失值分析(Y)...
多重归因(T)
复杂抽样(L)
质量控制(Q)
ROC 曲线图(V)...

被引次数

机构	被引次数
1 Zhejiang University	5520
2 Shanghai Jiao Tong University	3312
3 Peking University	8393
4 Tsinghua University	4799
5 Fudan University	2493
6 Sun Yat Sen University	6977
7 Sichuan University	5649
8 Huazhong University of Science & Technology	398429
9 Nanjing University	
10 Shandong University	
11 Jilin University	
12 University of Science & Technology of China	
13 Harbin Institute of Technology	
14 Xi'an Jiaotong University	
15 Central South University	
16 Wuhan University	
17 Tongji University	
18 Dalian University of Technology	
19 Southeast University - China	
20 Tianjin University	
21 South China University of Technology	
22 Beihang University	

两步聚类(T)...
K-均值聚类(K)...
系统聚类(H)...
树(R)...
判别(D)...
最近邻元素(N)...

数据视图 变量视图

系统聚类(H)...

IBM SPSS Statistics Processor 就绪

系统聚类分析

变量:

- 论文篇数
- 被引次数
- 篇均被引次数
- 学科规范化引文印...
- 高被引率

统计量(S)...
绘制(T)...
方法(M)...
保存(A)...

标注个案(C):

分群

☐ 个案 ☒ 变量

输出

☒ 统计量 ☒ 图

确定 粘贴(P) 重置(R) 取消 帮助

系统聚类分析: 统计量

×

☒ 合并进程表(A)

☐ 相似性矩阵(P)

聚类成员

☒ 无(N)

☐ 单一方案(S)

聚类数(B):

☒ 方案范围(R)

最小聚类数(M):

最大聚类数(X):

继续

取消

帮助

系统聚类分析: 图

×

☒ 树状图(D)

冰柱

☒ 所有聚类(A)

☐ 聚类的指定全距(S)

开始聚类(T):

停止聚类(P):

排序标准(B):

☒ 无(N)

方向

☒ 垂直(V)

☐ 水平(H)

继续

取消

帮助

系统聚类分析: 方法

×

聚类方法(M): 组间联接

度量标准

☒ 区间(N)

Euclidean 距离

幂: 根(R):

☐ 计数(I): 卡方度量

☐ 二分类(B): 平方 Euclidean 距离

存在(P): 不存在(A):

转换值

标准化(S): Z 得分

☒ 按照变量(V)

☐ 按个案:

转换度量

☐ 绝对值(L)

☐ 更改符号(H)

☐ 重新标度到 0-1 全距(E)

继续

取消

帮助

3.5-因子分析

□做一个形象的比喻，在观看电影时或观看以后，我们能够说出电影是否精彩，这是判别分析；并且我们会迅速地将对电影的印象形成两类：精彩或不精彩，把现在看的这部归入到某一类中，这是聚类分析；我们之所以可以认为这部电影精彩，是因为它具有精彩的影视作品所具有的一些共同特点：演员的演技、画面制作精良、讲述的故事有趣，等等。这种从研究对象中寻找公共因子的办法就是因子分析。

3.5-因子分析

- 简单地说，因子分析就是将原始变量分解成几个公共因子，在每个公共因子上有载荷的体现。如果一些原始变量在同一个公共因子上都具有较高的载荷，那么则说明这些原始变量有共同的内在（公共因子）。
- 例19：根据例19.sav提供的指标，设计一个具有3个一级指标的指标评价体系。

3.5-因子分析

例19.sav [数据集1] - IBM SPSS Statistics 数据编辑器

文件(F) 编辑(E) 视图(V) 数据(D) 转换(T) **分析(A)** 直销(M) 图形(G) 实用程序(U) 窗口(W) 帮助

报告
描述统计
表(T)
比较均值(M)
一般线性模型(G)
广义线性模型
混合模型(X)
相关(C)
回归(R)
对数线性模型(O)
神经网络
分类(F)
降维
度量(S)
非参数检验(N)
预测(T)
生存函数(S)
多重响应(U)
缺失值分析(Y)...
多重归因(T)
复杂抽样(L)
质量控制(Q)
ROC 曲线图(V)...

可见：8 变量的 8

引[次数	被引[率	热点论文率	高被引论文率
6.05	79.25	.04	1.09
6.30	79.76	.05	1.52
7.20	79.32	.10	2.00
7.70	83.76	.07	2.09
6.92	78.61	.04	1.52
6.12	77.48	.03	1.16
5.16	77.05	.02	.76
5.64	81.50	.03	.92
6.12	81.88	.07	1.46
6.12	81.88	.04	1.75
6.12	81.88	.02	.78
6.12	81.88	.06	1.43
6.26	79.87	.06	1.34
6.25	81.40	.13	1.87
8.59	84.84	.04	2.38
5.66	79.43	.00	1.21
7.07	81.89	.09	1.55

因子分析(F)...
对应分析(C)...
最优尺度(O)...

数据视图 变量视图

因子分析(F)...

IBM SPSS Statistics Processor 就绪



因子分析: 抽取



方法(M): 主成份

分析

- ☒ 相关性矩阵(R)
- ☐ 协方差矩阵(V)

输出

- ☒ 未旋转的因子解(F)
- ☐ 碎石图(S)

抽取

- ☐ 基于特征值(E)

特征值大于(A):

1

- ☒ 因子的固定数量(N)

要提取的因子(T):

3

最大收敛性迭代次数(X): 25

继续

取消

帮助

3.5-因子分析

结果显示

成份矩阵^a

	成份		
	1	2	3
论文数量	.508	.854	.087
被引次数	.662	.743	.034
篇均被引次数	.943	-.198	-.069
被引率	.714	-.160	-.539
热点论文率	.573	-.230	.721
高被引论文率	.837	-.365	.224
前10%论文率	.916	-.204	-.238

总结

- 理解数理统计的基本工具方法是关键
- 对输入数据的要求和对处理结果的分析解释是使用SPSS的主要工作
- 建议在掌握SPSS的同时学会一门程序语言

□讲座PPT会上传至讲座预约系统，可在系统中下载

第1讲 挖掘学习、研究的宝库——如何利用图书馆

分场安排:	时间	地点	适用对象	主讲教师	邮箱	预约人数
	3月22日 (周四) 16:00	工学图书馆	工科高年级本科生、研究生	梁金萍 	jpliang@scu.edu.cn	9 人预约
	3月28日 (周三) 16:00	文理图书馆	文理科高年级本科生、研究生	王圣洁 	wangshengjie@scu.edu.cn	19 人预约
	4月3日 (周二) 16:00	医学图书馆	医科高年级本科生、研究生	赵萍	zhaoping@scu.edu.cn	8 人预约

主要内容: 本讲座将分别介绍老校区文理、工学、医学3个分馆的概况、馆藏分布、常用电子资源的检索、书目查询方法，以及馆际互借、学科馆员、科技查新等特色服务，其目的在校区图书馆，充分利用图书馆资源。

□ 课件部分内容及数据来自《SPSS统计分析大全》（武松, 潘发明等编著，清华大学出版社）

- 索书号C819/1348，文理馆
- 随书光盘（包括教学视频和源数据）可以下载。

第 1 条记录(共 1 条)	
标准格式	S·F·X
系统号- 图书	000956148
ISBN	●978-7-302-34789-7 : CNY69.80 (含光盘) ●978-7-89414-781-3 (光盘)
作品语种	chi
题名	●SPSS统计分析大全 / 武松, 潘发明等编著
出版发行	●北京: 清华大学出版社, 2014
随书光盘	● 随书光盘下载
内容简介	《SPSS统计分析大全》由浅入深, 全面、系统地介绍了SPSS19.0的应用。《SPSS统计分析大全》涉及面广, 从软件基本操作到高级统计分析技术, 几乎涉及SPSS目前的绝大部分应用范畴。书中提供了大量应用案例, 供读者实战演练。另外, 《SPSS统计分析大全》配1张DVD光盘, 收录了作者为本书录制的16小时配套高清教学视频及书中所有案例的数据文件。《SPSS统计分析大全》共30章, 分为3篇。第1篇为 (更多)
目录:	第1篇 SPSS软件基础篇 第1章 SPSS19.0概述 1.1 SPSS19.0简介

