



Diapositivos de apoio às disciplinas de

Análise de Dados

Análise Avançada de Dados

Docente: Pedro Sousa (pedro.sousa.mrm@gmail.com)

Enquadramento

Estes diapositivos focam os pontos essenciais da matéria das duas disciplinas acima mencionadas – Análise de Dados (AD) e Análise Avançada de Dados (AAD), obedecendo aos conteúdos programáticos das mesmas. De modo algum dispensam que o aluno prossiga a sua investigação e estudo das matérias leccionadas, consultando para isso títulos constantes na lista bibliográfica cedida, ou outros.

1



Diapositivos de apoio à disciplina de

Análise de Dados

Docente: Pedro Sousa (pedro.sousa.mrm@gmail.com)

2

Análise de Dados 

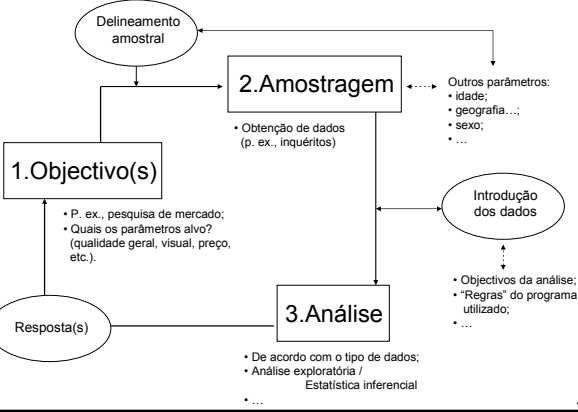
Tópicos preliminares

- Intro
- Horários das aulas
- Programa, bibliografia e avaliação
- Esquema de trabalho por tema (introdução, exemplos, ficha de trabalho)
- Material de apoio
- Registo no site INUAF (www.inuaf-studia.pt)
- Calendário
- Marcação de provas
- Delegado(s)
- Presenças
- Recomendações

3

Enquadramento geral...

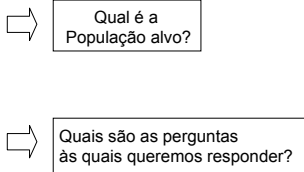
Análise de Dados



Objectivos da análise de dados

Análise de Dados

- Identificar e compreender o problema:

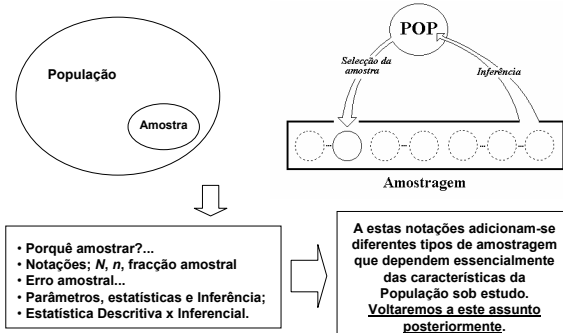


5

Objectivos da análise de dados

Análise de Dados

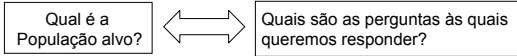
Intervalo estatístico: "3 MUNDOS"



6

Objectivos

- Identificar e compreender o problema:



- Tipos de fraldas descartáveis;*
- Intenção de voto para prox. AEINUAUF;*
- Vestuário...*
- Controle de qualidade de leitores de DVD;*
- Emissão de gases de combustão automóvel numa dada cidade.*
- ...

7

Objectivos

- Quais são as características da População que queremos estudar?



Variáveis aleatórias



- Traduzem a variação que existe na população.
- Obedecem a uma DISTRIBUIÇÃO DE PROBABILIDADES.



8

Objectivos

- A **TEORIA DAS PROBABILIDADES*** permite-nos **INFERIR!**



- "Qual é a verosimilhança** de observar ou amostrar um dado indivíduo ou um dado conjunto de indivíduos da População (amostra)?".



Ex. informal...:
Lotaria ou totoloto?...

9

O Programa SPSS

A progressiva utilização de computadores tem disponibilizado um número cada vez maior de métodos estatísticos que, embora já existissem, a sua aplicabilidade era restrita pela complexidade de cálculos.

Por outro lado, o tratamento de grandes quantidades de dados tem também vindo a ser praticável com o auxílio da informática.

Neste contexto, têm sido desenvolvidos vários programas informáticos de estatística, de entre os quais se destacam o **SPSS**, o **SAS**, o **Statistica**, o **St-PLUS** e o **R**. Este último tem a grande vantagem de ser gratuito, embora se exija que o utilizador possua conhecimentos médios de programação (mais informações sobre o **R** em <http://cran.r-project.org>).

De entre estes programas, o que utilizaremos na presente disciplina é o **SPSS**.

10

O Programa SPSS

Introdução

Consideremos uma amostra de 6 alunos que prestaram provas psicotécnicas, na qual se registou o *sexo*, *idade* e classificação (*nota*), de 0 a 100, na referida prova. A Tabela abaixo representa esta informação introduzida numa folha de cálculo (não no SPSS).

Sexo, *idade* e *nota* são variáveis. Os alunos são os indivíduos, casos ou observações.

	A	B	C	D
1	Nome	Sexo	Idade	Nota
2	Sara	f	16	87
3	Daniel	m	14	78
4	Pedro	m	17	98
5	Lucia	f	19	67
6	Sandra	f	21	50
7	Cristina	f	12	72

11

O Programa SPSS

Introdução

Quando abrimos o programa SPSS, ficamos presente o *Data Editor*, o qual possui dois menus: *Data View* (Figura à esquerda) e *Variable View* (Figura à direita). No *Data View*, como o nome indica, são introduzidos e visualizados os dados, tal como no presente exemplo (página anterior). No *Variable View*, são descritas as características das variáveis, tal como o nome, o tipo (numérico ou não), número de casas decimais (*Decimals*), a largura de coluna (*Columns*) e o alinhamento (*Align*). Os campos restantes serão explicados seguidamente.



12

O Programa SPSS

Introdução

Quando introduzimos os dados no SPSS tal qual o exemplo da folha de cálculo da página anterior, os dois menus anteriores assumem o aspecto abaixo.

Data View

	nome	sexo	idade	nota
1	Sara	f	16.00	87.00
2	Daniel	m	14.00	78.00
3	Pedro	m	17.00	98.00
4	Lucia	f	19.00	67.00
5	Sandra	f	21.00	50.00
6	Crstina	f	12.00	72.00

Repare-se que o SPSS adoptou duas casas decimais (*Decimals*) para as variáveis *idade* e *nota*. Podemos alterar para 0 casas decimais no campo *Decimals* do *Variable View*. O SPSS também reconheceu que *nome* e *sexo* não são variáveis numéricas, mas sim texto (*String*).

Variable View

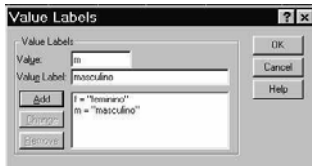
	Name	Type	Width	Decimals	Label	Values	Missing	Columns	Align	Measure
1	nome	String	8	0		None	None	8	Left	Nominal
2	sexo	String	8	0		(f, feminino)	None	8	Left	Nominal
3	idade	Numeric	11	2		None	None	8	Right	Ordinal
4	nota	Numeric	11	2		None	None	8	Right	Ordinal

13

O Programa SPSS

Introdução

O campo *Values* especifica o significado de siglas, letras ou outro tipo de caracteres que possam ser utilizados em variáveis não numéricas. Por exemplo, neste caso, "**f**" = **feminino** e "**m**"=**masculino**. O menu de introdução de *value labels* encontra-se na Figura abaixo.



NOTA IMPORTANTE: a variável "sexo" foi considerada, para este exemplo, como texto (*string*). No entanto, como veremos posteriormente, as variáveis categóricas (ou factores) possuem níveis, os quais assumem valores numéricos. Assim teríamos: *Value*: 1, *Value label*: masculino; *Value*: 2, *Value label*: feminino, por exemplo.

14

O Programa SPSS

Nomes de variáveis - REGRAS

Existem algumas regras para nomear variáveis (*Name*, menu *Variable View*), as quais se especificam seguidamente:

- Deve começar por uma letra ou por @ .
- Não pode terminar com ponto.
- Pode conter letras, dígitos ou qualquer dos seguintes caracteres: @, #, _ ou \$.
- Não pode conter espaços ou caracteres especiais tais como ! , ? e * , e outros que não os referidos acima .
- Não pode ser uma das palavras-chave do SPSS, ou seja:
AND, NOT, EQ, BY e ALL .

15

O Programa SPSS

Tipos de ficheiros SPSS / Gravar ficheiros.

Existem três tipos principais de ficheiros SPSS, nomeadamente, ficheiros de **dados** (*.sav), de **output** (*.spo) e de **syntax** (*.sps). Os ficheiros de **output** registam todos os procedimentos efectuados e respectivos resultados, em cada sessão de trabalho. O registo dos procedimentos é efectuado através de uma linguagem chamada **syntax**¹. Esta linguagem, e apenas esta, é utilizada nos ficheiros de **syntax**, e torna-se útil em muitos aspectos. Por exemplo, quando se pretende repetir uma série de procedimentos sem ter que utilizar os menus, ou quando se pretende programar qualquer operação não implementada directamente nos menus do SPSS. Veremos posteriormente casos concretos de aplicação da linguagem **syntax**.

Para gravar ficheiros efectua-se FILE \ SAVE. A Figura ao lado é um exemplo para o nosso ficheiro de dados.



¹ O processo de registar todos os procedimentos no **output** (display commands in the log) é opcional, sendo a sua activação / desactivação efectuada no sub-menu EDIT \ OPTIONS. Este sub-menu será explorado posteriormente.

O Programa SPSS

Ordenação de observações (sort)

Para ordenar (sort) observações no SPSS utiliza-se o sub-menu DATA \ SORT CASES. Por exemplo, ordenemos crescentemente (ascending) os nossos dados pela variável **nota** (Figura abaixo, à esquerda). O resultado no **Data View** encontra-se na Figura da direita.



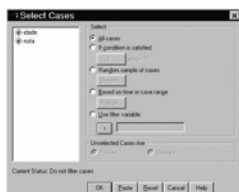
NOTA: Pode utilizar-se mais que uma condição no campo "sort by". Por exemplo, "nota-ascending" e "idade-descending". A primeira condição tem preferência relativamente à segunda. O mesmo se estende a mais condições.

	nome	sexo	idade	nota
1	Sandra	f	21	50
2	Lucia	f	19	67
3	Cristina	f	12	72
4	Carroll	m	14	78
5	Sara	f	16	87
6	Pedro	m	17	98

O Programa SPSS

Seleccção de observações (select cases)

Imaginemos que queremos analisar apenas uma parte do conjunto de dados, por exemplo, achar a média das notas para os indivíduos do sexo feminino. Para isso, utilizamos o menu DATA \ SELECT CASES.



Especifiquemos então que só queremos analisar os indivíduos do sexo feminino, no sub-menu "if condition is satisfied \ if". Note-se a variedade de hipóteses de operações de selecção disponíveis neste sub-menu.



O Programa SPSS

Seleção de observações (continuação)

O *Data View* passa então a apresentar o aspecto abaixo, no qual os indivíduos não seleccionados surgem "traçados" (excluídos da análise sub-sequente). Observe-se que o SPSS cria uma variável de *filtragem* (filter_\$), de natureza *binária* (0 ou 1), a qual utiliza na selecção de casos (0 = exclui; 1 = inclui). A elaboração de filtros deste género pode ser muito vantajosa, como veremos posteriormente.

Para voltar a analisar todos os casos novamente, executa-se
DATA \ SELECT CASES \ SELECT **All cases**.

	nome	sexo	idade	nota	filter_\$
1	Sara	f	16.00	87.00	1
2	Daniel	m	14.00	78.00	0
3	Pedro	m	17.00	98.00	0
4	Lucia	f	19.00	67.00	1
5	Sandra	f	21.00	50.00	1
6	Cristina	f	12.00	72.00	1

19

O Programa SPSS

Seleção de observações (continuação):

Filtros mais elaborados – ex.1

Imagine-se que se pretende seleccionar os indivíduos do sexo feminino **e** aqueles com idade igual ou superior a 17 anos. Nesta filtragem, bem como noutras de semelhante natureza, são utilizados os operadores lógicos:

"&" : "e", ou seja, observações que satisfaçam todas as condições explícitas;
"|" : "ou", pelo menos uma das condições.

Podem também ser utilizados todos os operadores matemáticos ou estatísticos disponíveis no SPSS, por exemplo:

"=" : "igual a"
"!=" : "diferente de"
">" : "maior que"
"<" : "menor que"
"<=" : "menor ou igual a"
">=" : "maior ou igual a"
"***" : "elevado a" (exponenciação)

NOTA: Para apagar variáveis ou casos (observações) **NÃO UTILIZAR** a tecla "delete" mas sim, com o **botão direito do rato**, seleccionar "**clear**".

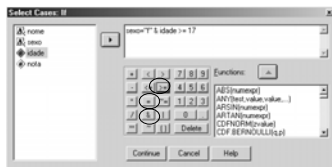
20

O Programa SPSS

Seleção de observações (continuação):

Filtros mais elaborados – ex.1

Assim, considerando as condições deste exemplo, temos:



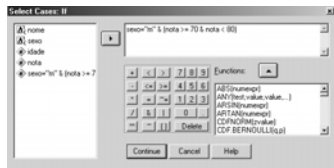
	nome	sexo	idade	nota	filter_\$
1	Sandra	f	21	50	Selected
2	Lucia	f	19	67	Selected
3	Cristina	f	12	72	Not Selected
4	Daniel	m	14	78	Not Selected
5	Sara	f	16	87	Not Selected
6	Pedro	m	17	98	Not Selected

21

O Programa SPSS

Seleção de observações (continuação):
Filtros mais elaborados – ex.2

Indivíduos masculinos cuja nota se situa entre 70 (inclusive) e 80:



	nome	sexo	idade	nota	filter_\$
1	Sandra	f	21	50	Not Selecte
2	Lucia	f	19	67	Not Selecte
3	Cristina	f	12	72	Not Selecte
4	Carroll	m	14	78	Selected
5	Sara	f	16	87	Not Selecte
6	Pedro	m	17	98	Not Selecte

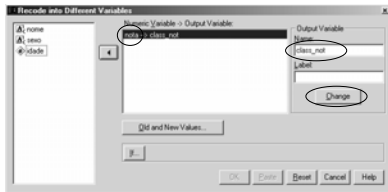
22

O Programa SPSS

Recodificar variáveis (recode)

Considere-se que se pretende classificar as notas em 3 categorias: 1 (50-69), 2 (70-89) e 3 (90-100) numa nova variável, por exemplo "class_not". Para isso utiliza-se o comando "TRANSFORM \ RECODE \ INTO DIFFERENT VARIABLES". Se se pretendesse substituir a variável original pela nova, codificada, seria utilizada a opção "... INTO SAME VARIABLES":

1. Seleccionar variável original;
2. Nomear output variable
3. Comando "change"
4. Comando "old and new values"



Passo seguinte

23

O Programa SPSS

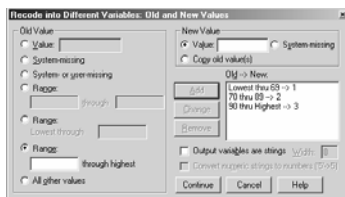
Recodificar variáveis (recode)

Especificar valores antigos e novas categorias:

Range: Lowest through 69
New Value / Value: 1
Add

Range: 70 through 89
New Value / Value: 2
Add

Range: 90 through highest
New Value / Value: 3
Add



"Do valor 90 até ao maior dos mesmos, atribui a categoria 3"

"Entre os valores 70 e 89 atribui a categoria 2"

"Do valor original mais baixo até ao 69, atribui a categoria 1"

Continue
OK

24

O Programa SPSS

Recodificar variáveis (recode)

Uma vez concluída a sequência de operações as alterações no *dataview* são as seguintes:

	nome	sexo	idade	nota	class not
1	Sandra	f	21	50	1.00
2	Lucia	f	19	67	1.00
3	Cristina	f	12	72	2.00
4	Daniel	m	14	78	2.00
5	Sara	f	16	87	2.00
6	Pedro	m	17	98	3.00
7					

25

O Programa SPSS

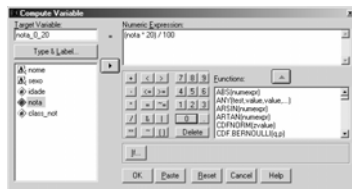
Cálculo de novas variáveis (compute)

Considere-se que se pretende converter a variável "nota", cuja amplitude é de 0 a 100, numa nova variável "nota_0_20", de amplitude 0 a 20 valores. A operação matemática consiste em:

$$\text{nota_0_20} = (\text{nota original} \times 20) / 100$$

No SPSS utiliza-se o comando "TRANSFORM \ COMPUTE":

1. Target variable
(nome da nova var.)
2. Numerical expression
(expressão de conversão)
3. OK



26

O Programa SPSS

Cálculo de novas variáveis (compute)

O resultado da operação encontra-se abaixo.

	nome	sexo	idade	nota	class not	nota_0_20
1	Sandra	f	21	50		10.00
2	Lucia	f	19	67		13.40
3	Cristina	f	12	72		14.40
4	Daniel	m	14	78		15.60
5	Sara	f	16	87		17.40
6	Pedro	m	17	98		19.60
7						

NOTA 1: Repare-se que é ainda possível especificar o tipo (*type*) e descrição (*label*) da nova variável, bem como especificar condições com o comando "if".

NOTA 2: O comando "paste", comum a todos os menus de operações em SPSS, faz com que todas as ordens dadas na operação sejam transcritas numa linguagem do SPSS (syntax) e adicionadas a uma janela à parte. Posteriormente estudaremos esta linguagem.

27

O Programa SPSS

Importação / exportação de ficheiros de dados

Tal como gravámos anteriormente o nosso ficheiro em formato *.SAV (SPSS), também é possível fazê-lo (**exportar o ficheiro**) em formato **Excel**. Para tal, no sub-menu FILE \ SAVE AS \ SAVE AS TYPE selecciona-se **Excel (*.XLS)**.

Para **importar** um ficheiro **Excel**, efectua-se FILE \ OPEN \ DATA \ FILES OF TYPE \ **Excel (*.XLS)**.

Na Figura ao lado encontra-se o exemplo de importação do conjunto de dados que tem vindo a ser trabalhado.

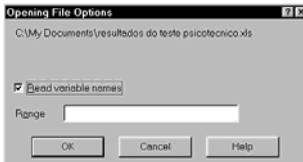


28

O Programa SPSS

Importação / exportação de ficheiros de dados

Após seleccionarmos **Open**, surge um novo menu que questiona se queremos que o SPSS leia o nome das variáveis (**read variable names**). Se na primeira linha do ficheiro *Excel* que estamos a importar se encontram os nomes das variáveis, esta opção deve ser seleccionada (Figura abaixo).



O resultado da importação deste ficheiro para o *Data View* do SPSS é apresentado na Figura do diapositivo 12.

29

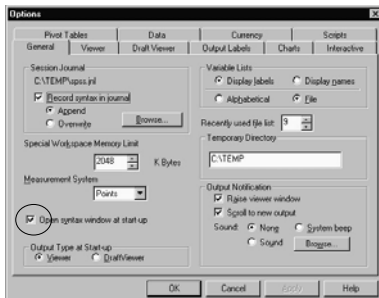
O Programa SPSS

Configuração do SPSS

Todas as opções de configuração do SPSS encontram-se no sub-menu EDIT \ OPTIONS, dentro do qual existem 10 outros sub-menus. Abordemos apenas dois deles, nomeadamente, EDIT \ OPTIONS \ GENERAL e EDIT \ OPTIONS \ VIEWER, e dentro destes, duas opções fundamentais.

No sub-menu GENERAL é importante assinalar a opção que automatiza a abertura de uma janela de *syntax* no arranque do SPSS:

Open syntax window at start-up (Figura ao lado).

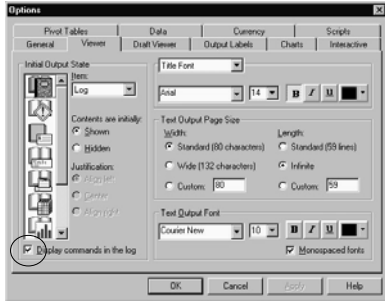


30

O Programa SPSS

Configuração do SPSS (continuação)

No sub-menu **Viewer** é fundamental assinalar a opção **Display commands in the log** que faz com que todos os procedimentos efectuados na sessão de trabalho sejam descritos no ficheiro de *output*, através da linguagem *syntax*.



31

Tipos de dados

O delineamento e a análise a efectuar depende do tipo de dados!

- **Qualitativos ou nominais** (categóricos, não-métricos)
 - Ordem subjacente → **ordinais**
 - **São categorias.**
Ex., fumador, sexo, profissão.
(casos part.: cat. binárias)
 - **São categorias ORDENADAS.**
Ex., critérios de preferência, grupos etários.
- **Quantitativos ou métricos**
 - **contínuos**
 - por intervalos **Escala contínua SEM ZERO ABSOLUTO.**
Ex., temperatura em °C.
 - por rácios **Escala contínua COM ZERO ABSOLUTO.**
Ex., tempo, altura, peso.
 - **discretos**
 - **Números inteiros.**
Ex., número de clientes por dia numa dada zona comercial.

32

Tipos de dados

Adaptado de um exemplo já trabalhado:

Trabalhador	Salário (€)	Categoria	Fumador	Agregado familiar
1	1024	2	1	3
2	567	1	1	4
3	892	1	0	2
4	1232	3	1	1
5	449	1	0	3
6	982	2	1	5
7	1100	2	0	2

- Quantitativos, discretos.
- Qualitativos, nominais (binários).
- Qualitativos, ordinais.
- Quantitativos, contínuos, por rácios.

Ex. SPSS...

33

A análise gráfica

- Porquê gráficos?
(uma imagem vale mais que mil palavras)

Mas não só!...



Caso de conjuntos
extensos de dados

(o que está directamente
relacionado com a elaboração
de tabelas de frequências, como
veremos posteriormente)

- Tipos de gráficos x tipos de dados!

34

A análise gráfica

- Pergunta:**
Com base na informação da tabela abaixo referente a preços (€)
de um dado produto em 500 diferentes zonas comerciais da
mesma cidade, identifique uma gama de preços mais praticados.
(sem recorrer ainda ao cálculo de indicadores estatísticos).

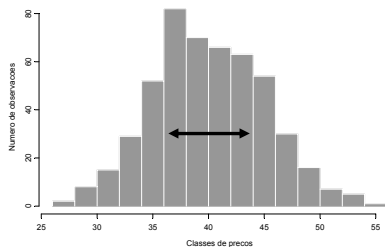
43.69	30.12	40.07	37.74	40.52	36.61	37.71	34.57	46.5	37.87	42.75	42.34	40.08	40.83	44.52	30.78	46.9	41.55	43.96	45.16	40.54	42.26
48.08	37.68	38.61	44.5	40.74	37.53	36.83	44.41	48.65	44.32	39.61	46.69	41.56	38.99	36.34	32.39	33.33	43.49	48.26	43.65	44.75	37.13
42.2	48.32	43.03	45.01	39.35	26.26	36.74	41.57	37.87	36.84	41.78	48.41	44	36.39	33.77	44.26	46.14	38.67	29.77	37.39	40.97	49.32
46.46	39.36	31.67	35.16	30.81	36.26	42.74	44.22	38.38	44.62	36.91	40.2	46.27	37.19	41.28	49.64	46.67	40.36	39.33	38.49	37.94	47.79
46.74	37.06	40.01	45.58	40.15	32.22	38.53	43.62	34.99	41.52	44.1	44.67	34.77	37.36	41.9	38.86	38.07	35.49	41.63	46.67	44.61	40.42
44.31	32.44	54.14	41.41	40.9	41.98	38.95	40.42	34.6	32.88	40.08	32.93	44.93	41.92	31.7	42.81	38.55	45.53	30.4	36.75	34.27	41.53
39.96	44.57	35.73	61.46	37.59	46.26	47	37.54	37.65	34.37	33.77	35.06	41.68	35.13	34.97	39.28	36.69	36.8	45.23	36.89	42.14	34.66
37.62	35.58	33.29	50.85	32.78	42.3	37.45	44.58	36.81	39.63	37.64	40	42.81	36.08	46.3	37.06	31.52	53.97	42.61	32.39	27.75	36.71
38.11	37.67	33.15	39.29	49.45	35.11	44.94	34.39	28.07	39.72	39.27	37.56	33.15	40.86	44.61	33.88	34.46	37.74	45.28	35.24	38.04	34.53
43.31	33.73	45.38	45.19	42.99	43.88	39.37	38.36	34.87	36.42	42.14	42.66	37.84	33.85	35.87	43.47	38.83	45.06	37.46	37.79	38.45	42.24
35.98	39.06	37.36	36.05	33.67	41.48	34.79	38.73	44.63	35.61	35.74	40.86	35.64	37.54	42.11	37.76	40.26	35.11	50.07	44.61	44.8	35.12
48.9	46.89	46.07	53.27	44.08	37.62	44.66	36.78	34.49	34.68	44.44	42.6	36.78	41.12	42.71	41.39	41.8	42.79	48.34	29.3	42.56	46.73
45.64	43.86	43.81	38.13	39.84	34.35	52.06	40.44	39.6	37.1	39.07	39.82	37.94	35.63	34.69	36.33	30.39	40.83	43.7	33.01	51.08	32.31
41.49	34.28	39.71	35.07	34.19	51.32	30.42	40.03	47.64	29.42	39.21	47.44	38.46	43.29	37.84	36.54	40.17	50.17	39.47	32.86	39.14	41.7
40.73	41.41	39.85	36.89	34.69	43.28	47.05	35.41	38.21	46.1	47.11	42.55	39.77	48.4	37.55	33.33	45.92	42	30.71	41.86	39.46	31.47
40.17	39.6	48.62	29.73	38.4	38.87	41.03	36.42	32.68	44.61	41.69	33.36	45.69	38.97	37.68	47.61	37.72	47.83	40.3	36.46	37.62	35.13
41.66	42.04	43.35	51.56	29.61	35	44.05	38.22	48.24	42.34	32.28	49.94	37.57	37.99	35.11	41.71	36.95	47.65	37.2	44.93	38.79	43.97
43.22	43.33	41.66	47.84	44.47	62.3	41.59	52.21	34.32	43.73	47.39	42.38	45.59	44.01	36.69	41.54	37.76	45.3	39.66	36.02	42.38	37.63
35.69	37.26	32.95	32.51	37.06	41.62	45.28	37.87	35.91	37.59	39.06	44.35	39.4	35.04	41.78	39.19	42.41	35.79	44.97	42.81	40.71	42.24
42.91	28.45	41.97	30.45	45.95	36.03	36.74	43.27	42.02	40.21	44.39	42.64	44.55	39.99	29.72	34.54	35.63	39.83	40.16	44.71	41.58	38.13

35

A análise gráfica

- Resposta:**

Histograma dos preços praticados



36

A análise gráfica

Resposta (cont.):

Na base da elaboração do **HISTOGRAMA** da página anterior está uma **TABELA DE FREQUÊNCIAS** na qual os dados originais são **agrupados em classes**:

1. Mínimo observado = 26.26 € ; máximo = 54.14 €.
2. Limite inferior da menor classe = 25 € ;
limite superior da maior classe = 55 €.
3. Estabelecer n° classes com base no intervalo de valores
(55 – 25) = 30. No exemplo, 15 classes. Assim, cada classe possui uma **AMPLITUDE** de (30 / 15) = 2 €.
4. Distribuir as observações pelas classes, obedecendo à regra:
≥ limite inferior da classe e < limite superior da classe*.
5. Os **CENTROS DE CLASSE***
surtem em consequência dos pontos anteriores.

37

A análise gráfica

TABELA DE FREQUÊNCIAS:

Classes		Frequências		
Limite inferior	Centro	absolutas	relativas (%)	rel. acumuladas (%)
25	26	1	0.2	0.2
27	28	3	0.6	0.8
29	30	15	3.0	3.8
31	32	20	4.0	7.8
33	34	40	8.0	15.8
35	36	56	11.2	27.0
37	38	86	17.2	44.2
39	40	67	13.4	57.6
41	42	73	14.6	72.2
43	44	64	12.8	85.0
45	46	34	6.8	91.8
47	48	22	4.4	96.2
49	50	9	1.8	98.0
51	52	7	1.4	99.4
53	54	3	0.6	100.0
total		500		

?:

- Frequências relativas
- Frequências acumuladas

38

A análise gráfica

Aplicaremos esta metodologia às variáveis do exemplo já trabalhado:

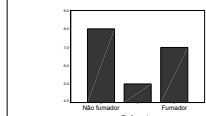
Trabalhador	Salário (€)	Categoria	Fumador	Agregado familiar
1	1024	2	1	3
2	567	1	1	4
3	892	1	0	2
4	1232	3	1	1
5	449	1	0	3
6	982	2	1	5
7	1100	2	0	2

39

Tipos básicos de gráficos de acordo com a natureza dos dados:

Dados qualitativos (nominais e ordinais) e dados quantitativos discretos

→ Gráficos de barras (bar charts)

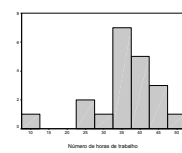


→ Gráficos "tipo tarte" (pie charts)



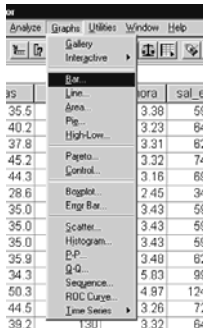
Dados quantitativos contínuos

Histogramas (histograms)



Elaboração de gráficos Em SPSS

Menu "Graphs" →



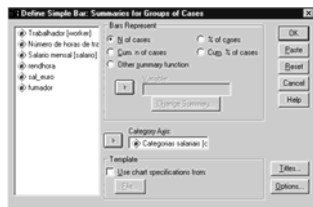
→ Cada sub-menu possui várias opções.

Analisemos, como exemplo, as relativas aos gráficos de barras.

Elaboração de gráficos Em SPSS

Gráficos de barras:
Variável "cat_sal" do exemplo já trabalhado.

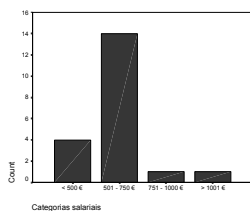
	Line	Area	Bar	Sal
35.5		3.38	56	
40.2		3.23	64	
37.8		3.31	62	
45.2		3.32	74	
44.3		3.16	66	
28.6		2.45	34	
35.0		3.43	56	
35.0		3.43	56	
35.9		3.48	62	
34.3		5.83	96	
50.3		4.97	124	
44.5		3.26	72	
39.2		3.32	64	



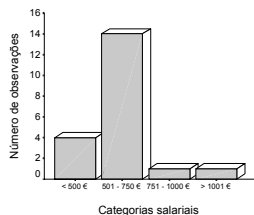
Elaboração de gráficos Em SPSS

Análise de Dados

Gráficos de barras
(*output*)



Gráficos de barras
(personalizado no editor de gráficos*)



Tipos de gráficos de barras x opções...

43

Elaboração de gráficos Em SPSS:

Análise de Dados

Gráficos de barras: *Stacked* (empilhado) e *clustered* (agrupado)

Após um debate televisivo no âmbito das eleições presidenciais nos Estados Unidos, recolheu-se, telefonicamente uma amostra $n=20$, tendo sido inquirida a intenção de voto e o nível de escolaridade do indivíduo. Os resultados encontram-se na tabela ao lado.

Pergunta: Graficamente será possível discernir alguma relação entre as duas variáveis?

	educação	voto_usa
1	Não completou ensino básico	Não vota
2	Ensino básico	Republicano
3	Ensino secundário	Democrata
4	Não completou ensino básico	Republicano
5	Não completou ensino básico	Não vota
6	Ensino básico	Republicano
7	Ensino secundário	Democrata
8	Não completou ensino básico	Não vota
9	Não completou ensino básico	Republicano
10	Não completou ensino básico	Democrata
11	Ensino básico	Republicano
12	Ensino secundário	Republicano
13	Ensino secundário	Democrata
14	Bacharelato	Democrata
15	Licenciatura ou superior	Democrata
16	Licenciatura ou superior	Democrata
17	Não completou ensino básico	Republicano
18	Ensino básico	Não vota
19	Ensino secundário	Republicano
20	Ensino secundário	Democrata

44

Elaboração de gráficos Em SPSS:

Análise de Dados

Gráficos de barras: *Stacked* (empilhado) e *clustered* (agrupado)

Resposta...

Graphs \
Bar \ *clustered* \
Define...

OK



NOTA: O sub-menu de elaboração de gráficos de barras empilhados (*stacked*) é semelhante ao apresentado ao lado (*clustered*)

45

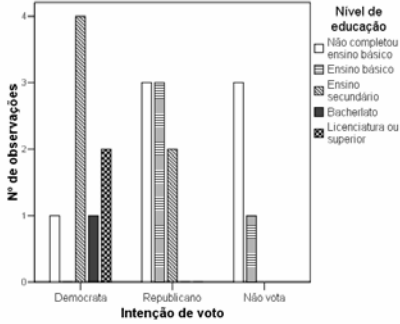
Elaboração de gráficos Em SPSS:

Análise de Dados

Gráficos de barras: *Stacked* (empilhado) e *clustered* (agrupado)

Resultados

(clustered):



46

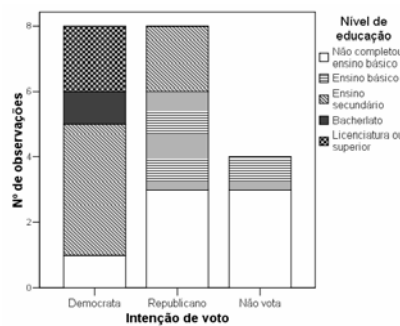
Elaboração de gráficos Em SPSS:

Análise de Dados

Gráficos de barras: *Stacked* (empilhado) e *clustered* (agrupado)

Resultados

(stacked):

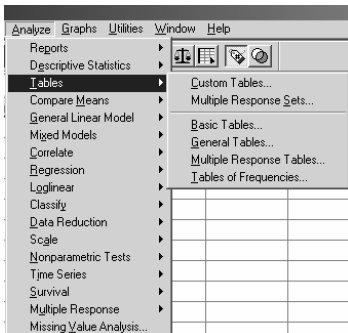


47

Elaboração de tabelas Em SPSS

Análise de Dados

Menu:
Analyze
Tables

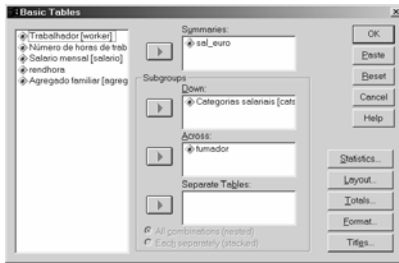


48

Elaboração de tabelas

Em SPSS: *Basic Tables*

Menu:
Analyze
Tables
Basic Tables

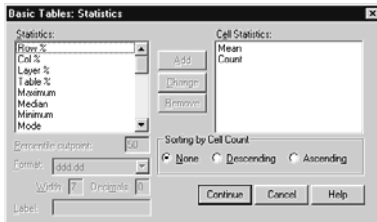


49

Elaboração de tabelas

Em SPSS: *Basic Tables* (continuação)

Menu:
Analyze
Tables
Basic Tables
Statistics



50

Elaboração de tabelas

Em SPSS: *Basic Tables* (output)

		Não fumador		Ex fumador		Fumador	
		n	Média	n	Média	n	Média
Categorias salariais	< 500 €	3	282.65			1	349.16
	501 até 750 €	5	628.49	4	623.50	5	660.41
	751 até 1000 €					1	997.60
	> 1000 €			1	1246.99		

51

Elaboração de tabelas

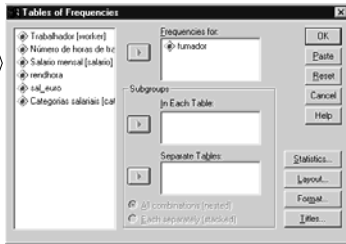
Em SPSS: Tabelas de frequências (tables of frequencies)

Menu:

Analyze

Tables

Tables of frequencies



52

Elaboração de tabelas

Em SPSS: Tabelas de frequências (tables of frequencies)

(continuação)

Menu:

Analyze

Tables

Tables of frequencies

Statistics



53

Elaboração de tabelas

Em SPSS: Tabelas de frequências (tables of frequencies)

(output)

	Observações	%
Não fumador	8	40.0%
Ex fumador	5	25.0%
Fumador	7	35.0%
Total	20	100.0%

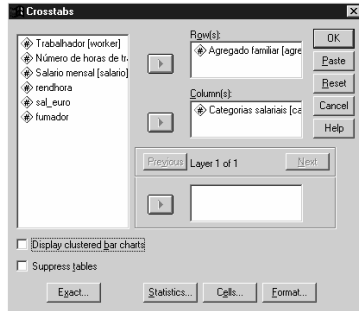
54

Elaboração de tabelas

Em SPSS: Tabelas cruzadas (*crosstabs*)

Menu:
Analyze
Descriptive Stats
Crosstabs

Weight cases...



55

Elaboração de tabelas

Em SPSS: Tabelas cruzadas (*crosstabs*)

(output)

Agregado familiar * Categorias salariais Crosstabulation

Count		Categorias salariais				Total
		< 500 €	501 até 750 €	751 até 1000 €	> 1000 €	
Agregado familiar	1		7			7
	2	1	2		1	4
	3	2	4			6
	4	1	1	1		3
Total		4	14	1	1	20

56

Análise exploratória

Indicadores de localização

$$\text{Média (mean)} \quad \bar{x} = \frac{1}{n} \sum_i x_i f_i \quad (\text{dados agrupados}) \quad \bar{x} = \frac{1}{n} \sum_i x_i \quad (\text{dados não agrupados})$$

$$\text{Mediana (median)} \quad M = \frac{N+1}{2} \quad (\text{para } N \text{ ímpar}) \quad M = \frac{x_{(N/2)} + x_{(N/2)+1}}{2} \quad (\text{para } N \text{ par})$$

Moda (*mode*): valor mais observado.

Quantis: Indicadores que têm a ver com a ordem das observações. São tanto mais úteis quanto maior for o conjunto de dados. Consistem em quartis (*quartiles*) quando se referem ao 1º, 2º e 3º quartos (25%, 50% e 75%) do número total de observações ordenadas crescentemente, ou em percentis (*percentiles*), quando se referem a qualquer outra localização referente aos valores ordenados. Por exemplo, o quantil₁₀ (percentil₁₀) refere-se à observação correspondente a ordem 10% do número total de observações ordenadas.

NOTA SOBRE: Observações atípicas: observações aberrantes relativamente ao conjunto de dados. A sua existência tanto pode ser devida a erros de introdução de dados, como pode traduzir um dado fenómeno real. Este facto sugere que sempre que existam observações atípicas se proceda a um estudo pormenorizado sobre a natureza das mesmas.

57

Análise exploratória

Indicadores de dispersão

Amplitude total (*range*): $A_{\text{at}} = x_n - x_1$

Amplitude inter-quartil: $A_{\text{iq}} = q_{75} - q_{25}$

Variação (*variance*):
$$s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$$

Desvio padrão (*standard deviation*):
$$s = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2}$$

Coefficiente de variação (*coefficient of variation*):

Consiste num indicador de dispersão relativa e é dado por: $CV = \frac{s}{\bar{x}}$

Outros indicadores

Assimetria (*skewness*): Tem a ver com a simetria da distribuição dos valores observados. Por exemplo, uma curva Normal possui assimetria 0. Valores de assimetria negativa indicam uma cauda da distribuição mais alongada para a esquerda. Valores positivos indicam o contrário.

Curtose (*kurtosis*): Tem a ver com o "peso das caudas" da distribuição e é comparado à curva Normal. Valores superiores a zero indicam caudas menos achatadas do que a Normal, e por isso de menor importância. Valores inferiores a zero indicam caudas mais achatadas do que a curva Normal, ou seja, de maior importância.

58

Análise exploratória

Representação gráfica de indicadores

Diagramas de caule-e-folhas (*stem-and-leaf plots*): Consistem em diagramas nos quais é representada uma categorização da variável em estudo, indicando os dígitos dominantes das observações (caules), os restantes dígitos das mesmas (folhas), a frequência do número de observações por categoria e ainda a presença de *outliers* ou de valores extremos.

Histogramas (*histograms*): São gráficos que categorizam a variável em estudo e nos quais a área de cada barra é proporcional à frequência de observações dentro da respectiva gama de valores.

Diagramas de extremos e quartis ou caixas de bigodes (*boxplots*): Neste diagramas são simultaneamente representados indicadores de localização e de dispersão. Dentro de algumas variantes, os mais comuns apresentam os valores mínimos e máximo (limites dos bigodes), mediana (2º quartil), e 1º e 3º quartis (limites das caixas).

A Figura seguinte descreve um exemplo de caixa de bigodes.

59

Análise exploratória

Caixas de bigodes (*boxplots*) - continuação:

Na figura ao lado apresenta-se um *boxplot* hipotético no qual:

A_{iq} : Amplitude inter-quartil (limites da caixa).

A_{ts} : Amplitude total com excepção de *outliers* e extremos (limites dos bigodes).

Repare-se que existem dois tipos de barreiras (a tracejado) que se estendem a partir dos limites da caixa, com base na Amplitude inter-quartil.

As **barreiras interiores** são definidas como:

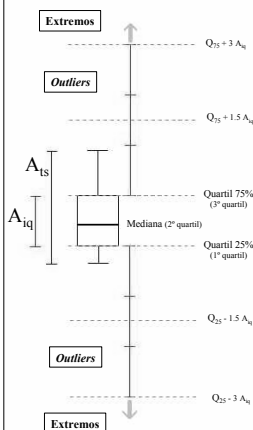
$$Q_{75} + 1.5 A_{\text{iq}} \\ Q_{25} - 1.5 A_{\text{iq}}$$

As **barreiras exteriores** são definidas como:

$$Q_{75} + 3 A_{\text{iq}} \\ Q_{25} - 3 A_{\text{iq}}$$

Assim, **outliers** são observações que se situam entre as barreiras interiores e as exteriores e **extremos** são observações que se situam além das barreiras exteriores.

60



Análise exploratória

Análise de Dados

Procedimentos em SPSS

EXEMPLO: A seguinte amostra traduz as quantidades (kg) de cana de açúcar cortada num dia por 15 trabalhadores num país da América latina:

140, 141, 145, 132, 112, 140, 141, 122, 111, 137, 132, 160, 128, 144, 129

Indicadores estatísticos

Analyze
Descriptive Statistics
Frequencies
Statistics Mean Median...kurtosis (indicadores)



61

Análise exploratória

Análise de Dados

Procedimentos em SPSS

Indicadores estatísticos (**output**):

Statistics		
Cana de açúcar cortada (kg)		
N	Valid	15
	Missing	0
Mean		134.27
Median		137.00
Mode		132 ^a
Std. Deviation		12.820
Variance		164.352
Minimum		111
Maximum		160
Sum		2014
Percentiles	25	128.00
	50	137.00
	75	141.00

a. Multiple modes exist. The smallest value is shown.

62

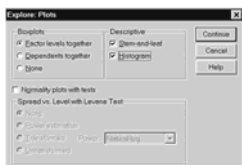
Análise exploratória

Análise de Dados

Procedimentos em SPSS

Representação gráfica de indicadores:

- diagramas de caule-e-folhas.
- histogramas.
- caixas de bigodes.



Analyze
Descriptive Statistics
Explore
Plots stem-and-leaf (caule-e-folhas)
boxplots (caixas de bigodes)
histogram

(**output**)

63

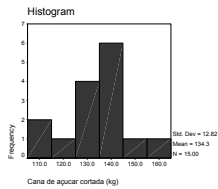
Análise exploratória

Análise de Dados

Procedimentos em SPSS

Representação gráfica de indicadores (output):

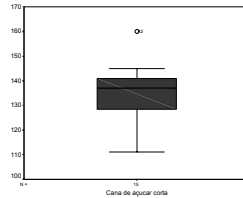
- diagramas de caule-e-folhas
- histogramas
- caixas de bigodes



Cana de açúcar cortada (kg) Stem-and-Leaf Plot

Frequency	Stem	Leaf
2	11	12
3	12	289
3	13	227
6	14	001145
1	15	001145
1	16	001145

Stem width: 10
Each leaf: 1 case(s)



64

Análise exploratória

Análise de Dados

Noção de relação entre duas variáveis métricas

Consideremos o seguinte exemplo: à medida que a idade de uma criança aumenta, a sua estatura aumenta também. Diz-se que há uma relação entre as variáveis "idade" e "estatura". Poderá dizer-se ainda que uma variável ("estatura") é função da outra ("idade"). Neste caso, a "estatura" que é função da "idade", isto é, que depende da "idade", denomina-se variável **dependente**. Por seu lado, a "idade", denomina-se variável **independente**.

Representação gráfica da relação entre duas variáveis métricas

Um dos primeiros passos no estudo da relação entre duas variáveis é elaborar uma **nuvem de pontos (scatterplot)**, isto é, um gráfico que traduza a variação da variável dependente em função da variação da variável independente. Seguindo o nosso exemplo anterior, consideremos uma amostra de 8 crianças e adolescentes dos quais se registou a idade e a estatura.

(exemplo)

65

Análise exploratória

Análise de Dados

Elaboração de uma nuvem de pontos (scatterplot)

Exemplo: consideremos uma amostra de 8 crianças e adolescentes dos quais se registou a idade e a estatura (cm):

	crianças	idade	estatura
1	1	14	155.00
2	2	6	120.00
3	3	11	145.00
4	4	7	123.00
5	5	13	161.00
6	6	12	150.00
7	7	2	91.00
8	8	4	103.00

⇒ Nuvem de pontos (SPSS):

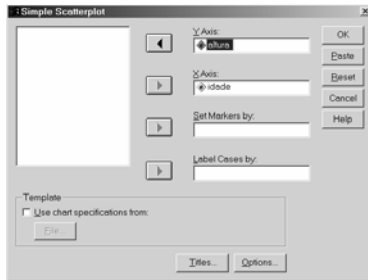
66

Análise exploratória

Análise de Dados

Nuvem de pontos (SPSS):

Graphs
Scatter...
Simple (define)
x axis
y axis



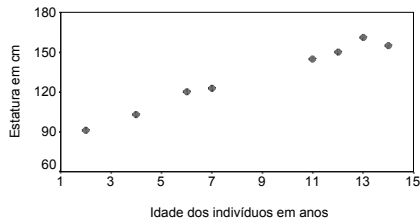
(output)

67

Análise exploratória

Análise de Dados

Nuvem de pontos (output):



Através da análise da nuvem de pontos podemos observar que, como esperaríamos, a "estatura" tende claramente a aumentar com a "idade", para os indivíduos amostrados. Neste caso, dizemos que parece existir uma **relação positiva** entre as duas variáveis. Noutros casos, a observação de uma dada nuvem de pontos pode sugerir que a relação entre duas variáveis é **nula**, ou que ela existe mas é **negativa**. Por outro lado, pode acontecer também que, embora seja evidente uma relação entre duas variáveis a mesma **não parece ser linear**. Exemplos...

68

Estatística Inferencial

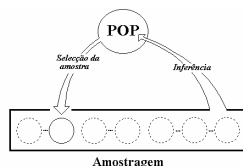
Análise de Dados

A análise exploratória permite APENAS conhecer a(s) amostra(s):

Respondemos a questões do género:

- Qual o formato da distribuição dos indivíduos da amostra?
- A maioria dos indivíduos amostrados são fumadores ou não-fumadores?
- Dos indivíduos amostrados, quem vai mais às compras, homens ou mulheres? Quais são os tipos de produtos comprados por uns e outros?
- ...

A Estatística Inferencial permite-nos ir mais além do universo amostral e tirar conclusões acerca da população da qual a amostra foi seleccionada.



69



Estatística Inferencial

Por exemplo, **com base numa amostra p** de n indivíduos extraída de uma população **P** de dimensão N ($n \ll N$) na qual se descrevem as variáveis sexo, dimensão do agregado familiar e rendimento mensal do agregado, podemos então responder a questões do género:

- Será que a proporção de homens na população P é 0.43?
- A média do número de indivíduos dos agregados familiares da população P é superior a 3?
- Será 890 € mensais o rendimento médio das famílias pertencentes à população P ?

70



Estatística Inferencial

Este tipo de suposições tem por base apenas uma população. Mas a Inferência pode estender-se a duas ou mais populações. Ambos os casos serão estudados posteriormente mas considerando duas populações P_1 e P_2 e duas amostras p_1 e p_2 extraídas respectivamente de P_1 e P_2 com as mesmas variáveis do exemplo acima, podemos então responder a questões que comparam efectivamente as pops:

- O rendimento médio dos agregados é diferente* nas pops P_1 e P_2 ?
- A proporção de homens em P_2 é superior a P_1 ?

* (des)igualdade matemática x estatística...

71



Estatística Inferencial

Os conceitos referidos abaixo são fundamentais em testes de hipóteses. Desenvolvamos cada um deles...

- Parâmetros (populacionais) e estatísticas (amostrais);
- Hipóteses nula (H_0) e alternativa (H_1);
- Nível de significância (α) de um teste;
- p -value (significância) do teste
- Testes unilaterais e bilaterais;
- Decisão (resposta formal) do teste;
- Tipos de erros (α e β);
- Pressupostos dos testes.

72

Estatística Inferencial

Os dois grandes grupos de testes: Paramétricos e não-paramétricos.

De uma forma breve (iremos abordar este assunto posteriormente) a diferença entre testes paramétricos e não-paramétricos tem a ver com PRESSUPOSTOS.

A aplicabilidade dos **testes paramétricos**, mais convencionais, dependem de pressupostos tais como a **Normalidade das observações***, entre outros. Muitas vezes este facto constitui uma limitação à análise quando os pressupostos não se verificam.

Os testes **não-paramétricos** não exigem este tipo de pressupostos**. Por este motivo a aplicabilidade dos mesmos é facilitada, embora de um modo geral estes métodos sejam ligeiramente menos robustos*** que os paramétricos.

Utilizaremos posteriormente testes paramétricos e não-paramétricos em paralelo, de acordo com o tipo de análise a efectuar.

73

Estatística Inferencial

Testes de hipóteses sobre parâmetros

Breve introdução teórica

Média : Comparação de $\bar{x}$ a μ_0 Estatística do teste: $t = \frac{\bar{x} - \mu_0}{s / \sqrt{n}}$

Teste bilateral Hipóteses: $H_0: \mu = \mu_0$; $H_1: \mu \neq \mu_0$

Decisão: Rejeitamos H_0 a um nível de significância α , se $|t| \geq t_{(\alpha/2)(n-1)}$

Teste unilateral Hipóteses:

$H_0: \mu = \mu_0$; $H_1: \mu > \mu_0$ Decisão: Rejeitamos H_0 a um nível de significância α , se $t > t_{\alpha(n-1)}$

$H_0: \mu = \mu_0$; $H_1: \mu < \mu_0$ Decisão: Rejeitamos H_0 a um nível de significância α , se $t < -t_{\alpha(n-1)}$

74

Estatística Inferencial

Testes de hipóteses sobre parâmetros

Breve introdução teórica

Proporção : Comparação de $\hat{p}$ a p_0

Hipóteses: $H_0: p = p_0$; $H_1: p \neq p_0$

Estatística do teste: $u = \frac{\hat{p} - p_0}{\sqrt{\frac{p_0(1-p_0)}{n}}}$

Decisão: Rejeitamos H_0 a um nível de significância α , se $|u| \geq Z_{(\alpha/2)}$

75

Estatística Inferencial

Análise de Dados



Testes de hipóteses sobre parâmetros

Breve introdução teórica

Variância : Comparação de s^2 a σ_0^2

Hipóteses: $H_0: \sigma^2 = \sigma_0^2$; $H_1: \sigma^2 \neq \sigma_0^2$

Estatística do teste: $\chi^2 = \frac{(n-1)s^2}{\sigma_0^2}$

Decisão: Rejeitamos H_0 a um nível de significância α , se $|\chi^2| \geq \chi^2_{(\alpha/2)(n-1)}$

76

Estatística Inferencial

Análise de Dados



Intervalos de confiança* na estimação de parâmetros

Breve introdução teórica

Consideremos três casos: média, proporção e variância.
Média

$$IC = \left[\bar{x} \pm t_{(\alpha/2)(n-1)} \frac{s}{\sqrt{n}} \right]$$

Proporção

$$IC = \left[\hat{p} \pm z_{(\alpha/2)} \sqrt{\frac{\hat{p}(1-\hat{p})}{n}} \right]$$

Variância

$$IC = \left[\frac{(n-1)s^2}{\chi^2_{(\alpha/2)(n-1)}}, \frac{(n-1)s^2}{\chi^2_{1-(\alpha/2)(n-1)}} \right]$$

Notações:

IC = Intervalo de confiança

$t_{(\alpha/2)(n-1)}$ = valor crítico de ordem $(\alpha/2)$ de uma lei de Student com $(n-1)$ graus de liberdade.

α = nível de significância

$z_{(\alpha/2)}$ = valor crítico de ordem $(\alpha/2)$ de uma lei Normal reduzida.

$\chi^2_{1-(\alpha/2)(n-1)}$ e $\chi^2_{(\alpha/2)(n-1)}$ = valores críticos de ordens $1-(\alpha/2)$ e $(\alpha/2)$, respectivamente, de uma lei de Qui-quadrado com $(n-1)$ graus de liberdade.

77

Estatística Inferencial

Análise de Dados



Comparação (paramétrica) entre duas populações Breve introdução teórica

Médias :

Comparação de médias de duas amostras independentes

Hipóteses: $H_0: \mu_1 = \mu_0$; $H_1: \mu_1 \neq \mu_0$

$$\text{Estatística do teste: } t = \frac{\bar{x}_1 - \bar{x}_2}{s_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}}$$

$$\text{sendo: } s_p = \sqrt{\frac{(n_1-1)s_1^2 + (n_2-1)s_2^2}{n_1 + n_2 - 2}}$$

Decisão: Rejeitamos H_0 a um nível de significância α , se $|t| \geq t_{(\alpha/2)(n_1+n_2-2)}$

Amostras independentes / amostras emparelhadas

Considerando duas amostras (1 e 2), independência implica que cada observação da amostra 1 não possui nenhum tipo de relação com nenhuma observação da amostra 2.

Emparelhamento implica que cada observação da amostra 1 possui uma determinada relação com uma observação da amostra 2, de tal modo que se pode considerar que os dados ocorrem aos pares.

Exemplo: Níveis de colesterol em 10 pacientes antes (amostra 1) e depois (amostra 2) de um programa de dieta e exercício. Note-se que os indivíduos das amostras 1 e 2 são os mesmos. Por esta razão, existe um emparelhamento.

78

Estatística Inferencial

Análise de Dados



Comparação (paramétrica) entre duas populações

[Breve introdução teórica](#)

Médias :

Comparação de médias de duas amostras emparelhadas

A ideia fundamental no tratamento de amostra emparelhadas é substituir o conjunto das duas amostras por uma única amostra, a amostra das diferenças dos valores de cada par. A partir daqui podemos recorrer aos métodos de inferência sobre parâmetros de localização, entre os quais, ao teste t , já referido, aplicado neste caso para inferir sobre a média populacional das diferenças entre as duas amostras.

Assim, tendo $\bar{X}_d$ = média da amostra das diferenças, queremos comparar $\bar{X}_d$ a μ_0 .

Re-escrevamos então as hipóteses nula e alternativa do seguinte modo:

$H_0: \mu_d = \mu_0$; $H_1: \mu_d \neq \mu_0$, sendo $\mu_d = \mu_1 - \mu_2$

Estatística do teste:

$$t = \frac{\bar{X}_d - \mu_0}{s_d / \sqrt{n}}$$

Decisão: Rejeitamos H_0 a um nível de significância α ,

se $|t| \geq t_{(\alpha/2)(n_1+n_2-2)}$

79

Estatística Inferencial

Análise de Dados



Comparação (paramétrica) entre duas populações

[Breve introdução teórica](#)

Variâncias

Hipóteses: $H_0: \sigma_1^2 = \sigma_2^2$; $H_1: \sigma_1^2 \neq \sigma_2^2$

Estatística do teste: $F = \frac{s_1^2}{s_2^2}$ isto é, o quociente das variâncias amostrais. Na prática, pode inverter-se o numerador / denominador de forma a que o quociente seja >1 (variância do numerador maior). Isto porque se simplifica a decisão do teste:

Decisão: Rejeitamos H_0 a um nível de significância α , se $F \geq F_{(\alpha/2)(n_1-1)(n_2-1)}$

Nota: Se o quociente entre as duas variâncias amostrais fosse <1 (variância do numerador menor que a do denominador), a região crítica do teste seria $F_{1-(\alpha/2)(n_1-1)(n_2-1)}$

80

Estatística Inferencial

Análise de Dados



Procedimentos em SPSS

[Duas advertências sobre testes de hipóteses em SPSS:](#)

1. REGRAS DE DECISÃO DO TESTE: **p-values**

Até aqui falámos no cálculo de estatísticas e na sua comparação com valores críticos tabelados, de forma a suportar a decisão do nosso teste. O SPSS, tal como a maioria dos outros programas de estatística, não nos fala de valores tabelados, mas sim dá-nos directamente o **p-value**, ou seja, a probabilidade exacta de, sendo H_0 verdadeira, obtermos um valor de teste igual ao que calculámos. Assim, num teste **bilateral**, se o nosso nível de significância (α) for de 0.05, rejeitamos H_0 se obtermos um **p-value** inferior a 0.05. Por outro lado, num teste **unilateral**, e considerando o mesmo valor de α , rejeitamos H_0 se o **quociente (p-value / 2)** for inferior a 0.05.

No caso onde pode haver bilateralidade ou unilateralidade do teste estatístico (caso de comparação entre valores amostrais a valores padrão, ou caso de comparação entre duas amostras), as significâncias calculadas pelo SPSS são sempre bilaterais - **Sig. (2-tailed)** ou **Asymp. Sig. (2-tailed)**. Se o teste elaborado pelo utilizador possuir carácter unilateral ter-se-á que proceder à transformação acima referida.

Consideremos o exemplo seguinte:

81

Estatística Inferencial

Análise de Dados

Procedimentos em SPSS

Duas advertências prévias sobre testes de hipóteses em SPSS:

1. REGRAS DE DECISÃO DO TESTE: *p-values* (continuação)

EXEMPLO (*): Num teste *t* de comparação entre duas médias amostrais, o valor de *t* calculado foi 2.718. O valor de α considerado foi 0.05, e o tamanho da amostra $n = 12$. A significância bilateral do teste (*p-value* ou *sig. (2-tailed)*) calculada pelo SPSS é 0.02. Assim, num teste bilateral, rejeitamos H_0 visto o valor da significância ser inferior a α . Se o teste for unilateral, a respectiva significância seria $0.02/2 = 0.01$. Também neste caso se rejeita H_0 uma vez que o valor da significância unilateral do teste é inferior a α .

2. Estatísticas e testes que o SPSS não possui directamente

Na realidade o SPSS não permite calcular directamente intervalos de confiança para proporções e para variâncias, testes paramétricos sobre proporções e testes a variâncias de uma ou duas amostras. No entanto, existe uma alternativa não-paramétrica para testes a proporções (teste binomial), e também a possibilidade de calcularmos as estatísticas e significâncias acima referidas através de funções em *syntax*, como veremos posteriormente.

(*) – Este exemplo foi elaborado com base numa tabela de valores críticos da distribuição *t*.

82

Estatística Inferencial

Análise de Dados

Procedimentos em SPSS

Particularidade na introdução de dados:

No caso de comparação entre **duas amostras independentes**, os valores das duas amostras terão que ser introduzidos numa mesma variável (variável de teste). Uma segunda variável especificará a pertença de cada valor à primeira e à segunda amostra (ex. ao lado).

	teste	amostras
1	14	1
2	10	1
3	10	1
4	10	1
5	10	1
6	13	2
7	10	2
8	14	2
9	14	2
10	12	2
11	14	2

83

Estatística Inferencial

Análise de Dados

Procedimentos em SPSS

Particularidade na introdução de dados:

No caso de comparação entre **duas amostras emparelhadas**, as amostras são introduzidas em variáveis separadas (ex. ao lado).

	amost1	amost2
1	38.8	26.8
2	51.3	34.8
3	43.8	31.8
4	64.9	56.6
5	29.8	29.0
6	43.4	37.2
7	44.8	36.3
8	33.9	25.2
9	62.7	42.2
10	40.1	29.3
11	55.3	44.1
12	47.4	46.1

84

Estadística Inferencial

Análise de Dados



Procedimentos em SPSS

Intervalos de confiança para médias	Analyze Descriptive statistics Explore Statistics Descriptives confidence interval for mean
Teste à média de uma população	Analyze Compare means One-sample t test nome da variável / valor a testar
Teste à diferença entre as médias de duas populações com base em amostras independentes	Analyze Compare means Independent samples t test test variable variável de estudo grouping variable variável de agrupamento group 1 (código 1) group 2 (código 2)
Teste à diferença entre as médias de duas populações com base em amostras emparelhadas	Analyze Compare means Paired sample t test variável 1 variável 2

85

Estadística Inferencial

Análise de Dados



Procedimentos em SPSS

Cálculo de estatísticas de teste e de p-values em SYNTAX

Como referido anteriormente, quando o SPSS não dispõe de caixas de diálogo para um determinado teste, pode recorrer-se à linguagem de syntax para calcular a estatística do teste e o *p-value*, de forma a decidirmos sobre a rejeição ou não da nossa H_0 .

Exemplo: comparação de $\hat{p} = 0.66$ a $P_0 = 0.5$, com base numa amostra $n = 20$.

Na janela de syntax:

<pre>matrix. compute u=(0.66-0.5)/sqrt((0.5*(1-0.5)/20)). compute p=1-cdfnorm(u). print u. print p. end matrix.</pre>	início do comando "matrix" cálculo de u de acordo com a exp. da pag 3 cálculo do p-value (p) imprimir o valor de u no output imprimir o valor de p no output fim do comando "matrix".
---	--

<u>Resultados (janela de output)</u> <pre>Run MATRIX procedure: U 1.431083506 P 10 ** -2 X 7.620314198 ----- END MATRIX -----</pre>	<u>COMENTÁRIOS:</u> "sqrt" = raiz quadrada; "cdfnorm" = função do SPSS que calcula a função distribuição acumulada da Normal reduzida <u>até</u> ao valor que lhe damos, neste caso, $u=1.431$. Como o que queremos não é o acumulado mas a área "restante", isto é, a área de decisão, calculamos $1 - \text{cdfnorm}(u)$ para o cálculo do <i>p-value</i>. Neste caso, $p=0.076$. A decisão do teste bilateral a $\alpha=0.05$ seria não rejeitar H_0, pois o <i>p-value</i> não é inferior a 0.05.
--	---

86

Estadística Inferencial

Análise de Dados



Procedimentos em SPSS

Cálculo de estatísticas de teste e de p-values em SYNTAX

Outras funções importantes no cálculo de *p-values* em syntax:

CHICDF(m,v)	Função distribuição de qui-quadrado. Com base na estatística por nós calculada (m) e dos graus de liberdade (v) a função dá-nos o <i>p-value</i> . <u>Exemplos de aplicação:</u> estimação de intervalos de confiança para variâncias, teste de comparação de uma variância a um valor padrão.
FCDF(m,v1,v2)	Função distribuição de F. Do mesmo modo: "m"=estatística do teste; "v1"=graus de liberdade do numerador; "v2"=graus de liberdade do denominador. <u>Exemplos de aplicação:</u> teste de comparação de duas variâncias.
TCDF(m,v)	Função distribuição de t - Student. "m"=estatística do teste; "v"=graus de liberdade. <u>Exemplos de aplicação:</u> Cálculo imediato de significâncias do valor testado.

87

Estatística Inferencial N-P

Análise de Dados



Comparação entre estatística paramétrica e não-paramétrica

Já abordámos previamente o problema da comparação de valores de uma determinada variável em duas populações diferentes, com o objectivo de decidir se é ou não admissível a hipótese de a generalidade dos valores das 2 populações serem globalmente análogos. O ponto de partida para esse estudo são duas amostras (uma de cada população), que podem ser extraídas de forma independente ou através de um delineamento que garanta o seu emparelhamento. A comparação das populações fazia-se mediante a comparação dos indicadores de localização média populacional, concretamente, através de testes à diferença das médias das duas populações.

No entanto, a validade dos testes estudados anteriormente depende principalmente da hipótese de normalidade da distribuição dos valores (ou das diferenças, no caso de amostras emparelhadas) populacionais. Embora os resultados relativos a médias sejam relativamente robustos a desvios à normalidade, têm vindo a ser cada vez mais utilizados os métodos não-paramétricos para um conjunto de testes cuja validade não depende da normalidade dos indivíduos da população.

Seguidamente serão estudados dois desses testes que são utilizados na comparação da localização de populações sem a hipótese de normalidade: **o teste de Wilcoxon-Mann-Whitney** (*Rank sum test*) válido no caso de amostras independentes e **o teste de Wilcoxon** (*Signed rank test*), válido no caso de amostras emparelhadas. Ambos os testes partilham a particularidade de substituírem os valores observados pela ordem (*rank*) dos mesmos.

88

Estatística Inferencial N-P

Análise de Dados



O teste de Mann-Whitney* (*rank sum test*)

Consideremos duas amostras independentes ordenadas por ordem crescente (ou decrescente), sendo n_1 = número de observações da amostra 1 e n_2 = número de observações da amostra 2. Suponhamos que, para efeitos de generalização, $n_1 \neq n_2$. A menor das observações do conjunto das duas amostras atribui-se a ordem 1, à segunda menor observação atribui-se a ordem 2, e assim consecutivamente até os $N = n_1 + n_2$ indivíduos se encontrarem ordenados. Note-se que, na ordenação não se perde a correspondência entre o indivíduo e a sua ordem respectiva.

É calculada então a estatística U de WMW:
$$U = n_1 n_2 + \frac{n_1(n_1 + 1)}{2} - R_1$$

Sendo R_1 a soma das ordens da amostra 1. No caso do teste ser bilateral é necessário também calcular a estatística U' . O maior valor entre U e U' é então comparado com o valor crítico $U_{(2), (n_1, n_2)}$. U' é então dado por:

$$U' = n_1 n_2 - U$$

* Este teste foi primeiramente elaborado por Wilcoxon e posteriormente desenvolvido por Mann e Whitney. No entanto, na bibliografia e programas de estatística encontra-se referido, na maioria das vezes, apenas por Mann-Whitney.

89

Estatística Inferencial N-P

Análise de Dados



O teste de Mann-Whitney (*rank sum test*)

Exemplo*: Queremos testar, a um nível de significância $\alpha=0.05$, a hipótese bilateral de não existir diferenças entre a altura de estudantes do sexo masculinos e feminino.

H_0 : Alunos e alunas são da mesma altura
 H_1 : Alunos e alunas não são da mesma altura.

Altura do alunos	Altura das alunas	Ordens das Alturas do alunos	Ordens da Alturas das alunas
170	163	4	1
178	165	7	2
180	168	8	3
183	173	9	5
185	175	10	6
188		11	
193		12	
$n_1=7$	$n_2=5$	$R_1=61$	$R_2=17$

$$U = 7*5 + ((7*(7+1))/2) - 61$$
$$U = 2$$

$$U' = 7*5 - 2$$
$$U' = 33$$

*Adaptado de Zar (1984).

90

Estatística Inferencial N-P

Análise de Dados

O teste de Mann-Whitney (*rank sum test*)

Exemplo (continuação): Consultemos então uma tabela de valores críticos da distribuição U de WMW, para um nível de significância de 0.05, $n_1 = 5$, $n_2 = 7$:

TABLE B.10 (cont.) CRITICAL VALUES OF THE MANN-WHITNEY U DISTRIBUTION

$n_1 \backslash n_2$	0.10	0.15	0.20	0.25	0.30	0.35	0.40	0.45	0.50
5 \ 7	91	98	104	110	116	122	128	134	140
5 \ 8	96	103	109	115	121	127	133	139	145
5 \ 9	102	109	115	121	127	133	139	145	151
5 \ 10	107	114	120	126	132	138	144	150	156
5 \ 11	112	119	125	131	137	143	149	155	161
5 \ 12	117	124	130	136	142	148	154	160	166
5 \ 13	122	129	135	141	147	153	159	165	171
5 \ 14	127	134	140	146	152	158	164	170	176
5 \ 15	132	139	145	151	157	163	169	175	181
5 \ 16	137	144	150	156	162	168	174	180	186
5 \ 17	142	149	155	161	167	173	179	185	191
5 \ 18	147	154	160	166	172	178	184	190	196
5 \ 19	152	159	165	171	177	183	189	195	201
5 \ 20	157	164	170	176	182	188	194	200	206
5 \ 21	162	169	175	181	187	193	199	205	211
5 \ 22	167	174	180	186	192	198	204	210	216
5 \ 23	172	179	185	191	197	203	209	215	221
5 \ 24	177	184	190	196	202	208	214	220	226
5 \ 25	182	189	195	201	207	213	219	225	231
5 \ 26	187	194	200	206	212	218	224	230	236
5 \ 27	192	199	205	211	217	223	229	235	241
5 \ 28	197	204	210	216	222	228	234	240	246
5 \ 29	202	209	215	221	227	233	239	245	251
5 \ 30	207	214	220	226	232	238	244	250	256
5 \ 31	212	219	225	231	237	243	249	255	261
5 \ 32	217	224	230	236	242	248	254	260	266
5 \ 33	222	229	235	241	247	253	259	265	271
5 \ 34	227	234	240	246	252	258	264	270	276
5 \ 35	232	239	245	251	257	263	269	275	281
5 \ 36	237	244	250	256	262	268	274	280	286
5 \ 37	242	249	255	261	267	273	279	285	291
5 \ 38	247	254	260	266	272	278	284	290	296
5 \ 39	252	259	265	271	277	283	289	295	301
5 \ 40	257	264	270	276	282	288	294	300	306
5 \ 41	262	269	275	281	287	293	299	305	311
5 \ 42	267	274	280	286	292	298	304	310	316
5 \ 43	272	279	285	291	297	303	309	315	321
5 \ 44	277	284	290	296	302	308	314	320	326
5 \ 45	282	289	295	301	307	313	319	325	331
5 \ 46	287	294	300	306	312	318	324	330	336
5 \ 47	292	299	305	311	317	323	329	335	341
5 \ 48	297	304	310	316	322	328	334	340	346
5 \ 49	302	309	315	321	327	333	339	345	351
5 \ 50	307	314	320	326	332	338	344	350	356
5 \ 51	312	319	325	331	337	343	349	355	361
5 \ 52	317	324	330	336	342	348	354	360	366
5 \ 53	322	329	335	341	347	353	359	365	371
5 \ 54	327	334	340	346	352	358	364	370	376
5 \ 55	332	339	345	351	357	363	369	375	381
5 \ 56	337	344	350	356	362	368	374	380	386
5 \ 57	342	349	355	361	367	373	379	385	391
5 \ 58	347	354	360	366	372	378	384	390	396
5 \ 59	352	359	365	371	377	383	389	395	401
5 \ 60	357	364	370	376	382	388	394	400	406
5 \ 61	362	369	375	381	387	393	399	405	411
5 \ 62	367	374	380	386	392	398	404	410	416
5 \ 63	372	379	385	391	397	403	409	415	421
5 \ 64	377	384	390	396	402	408	414	420	426
5 \ 65	382	389	395	401	407	413	419	425	431
5 \ 66	387	394	400	406	412	418	424	430	436
5 \ 67	392	399	405	411	417	423	429	435	441
5 \ 68	397	404	410	416	422	428	434	440	446
5 \ 69	402	409	415	421	427	433	439	445	451
5 \ 70	407	414	420	426	432	438	444	450	456
5 \ 71	412	419	425	431	437	443	449	455	461
5 \ 72	417	424	430	436	442	448	454	460	466
5 \ 73	422	429	435	441	447	453	459	465	471
5 \ 74	427	434	440	446	452	458	464	470	476
5 \ 75	432	439	445	451	457	463	469	475	481
5 \ 76	437	444	450	456	462	468	474	480	486
5 \ 77	442	449	455	461	467	473	479	485	491
5 \ 78	447	454	460	466	472	478	484	490	496
5 \ 79	452	459	465	471	477	483	489	495	501
5 \ 80	457	464	470	476	482	488	494	500	506
5 \ 81	462	469	475	481	487	493	499	505	511
5 \ 82	467	474	480	486	492	498	504	510	516
5 \ 83	472	479	485	491	497	503	509	515	521
5 \ 84	477	484	490	496	502	508	514	520	526
5 \ 85	482	489	495	501	507	513	519	525	531
5 \ 86	487	494	500	506	512	518	524	530	536
5 \ 87	492	499	505	511	517	523	529	535	541
5 \ 88	497	504	510	516	522	528	534	540	546
5 \ 89	502	509	515	521	527	533	539	545	551
5 \ 90	507	514	520	526	532	538	544	550	556
5 \ 91	512	519	525	531	537	543	549	555	561
5 \ 92	517	524	530	536	542	548	554	560	566
5 \ 93	522	529	535	541	547	553	559	565	571
5 \ 94	527	534	540	546	552	558	564	570	576
5 \ 95	532	539	545	551	557	563	569	575	581
5 \ 96	537	544	550	556	562	568	574	580	586
5 \ 97	542	549	555	561	567	573	579	585	591
5 \ 98	547	554	560	566	572	578	584	590	596
5 \ 99	552	559	565	571	577	583	589	595	601
5 \ 100	557	564	570	576	582	588	594	600	606

Decisão:

Uma vez que o valor crítico do teste é $U_{0.05(2),(5,7)} = 30$ (tabela), rejeitamos H_0 pois $U = 33 > 30$.

Concluimos que, a um nível de significância $\alpha = 0.05$, os alunos e alunas não são da mesma altura.

91

Estatística Inferencial N-P

Análise de Dados

O teste de Mann-Whitney (*rank sum test*)

Casos particulares:

1. A existência de "empates" (*tied ranks*)

No exemplo anterior não encontramos, no conjunto das duas amostras, observações iguais. No entanto, pode acontecer, e neste caso estamos em presença de "empates" (*tied ranks*). Nesse caso, as ordens de observações de igual valor são a média das ordens que teriam sido atribuídas a estas observações, caso não existisse empate. Como exemplo consideremos o exercício anterior. Se o 4º valor mais elevado das alunas fosse 170 e não 173, teríamos um empate com o valor mais baixo dos alunos. Concretamente, entre as ordens 4 e 5. Neste caso atribuir-se-ia a ordem $(4+5)/2 = 4.5$ a cada uma das observações. Uma vez efectuada esta alteração, o procedimento segue como referido anteriormente.

2. O teste unilateral

Nos testes unilaterais há que decidir se calcular U ou U' , de acordo com o tipo de hipótese a testar (quadro ao lado). O valor crítico tabelado será

$U_{\alpha(1),(n_1,n_2)}$

	H_0 : grupo 1 $\geq$ grupo 2 H_1 : grupo 1 < grupo 2	H_0 : grupo 1 $\leq$ grupo 2 H_1 : grupo 1 > grupo 2
Ordenação crescente	U	U'
Ordenação decrescente	U'	U

92

Estatística Inferencial N-P

Análise de Dados

O teste de Mann-Withney (*rank sum test*)

Procedimento em SPSS

Introdução dos dados

À semelhança do que já vimos para testes t a duas amostras independentes, uma variável contém os valores das duas amostra, havendo uma segunda variável que especifica os códigos de pertença de cada valor à variável 1 ou 2.

Menus

Analyze

Nonparametric Tests

2 Independent Samples

Test variable / grouping variable

Test Type Mann-Whitney U

OK

93

Estatística Inferencial N-P

O teste de Wilcoxon (*signed-rank test*)

Do mesmo modo que o teste de Wilcoxon-Mann-Whitney constitui uma alternativa não-paramétrica ao teste t para duas amostras independentes, o teste de Wilcoxon apresenta-se como alternativa não-paramétrica ao teste t para duas amostras emparelhadas.

O teste consiste nos seguintes procedimentos sequenciais:

1. Cálculo das diferenças entre os pares de observações (tal como o teste t para amostras emparelhadas).
2. Ordenação crescente dos valores absolutos das diferenças.
3. Atribuição do sinal de cada diferença (- ou +) aos valores ordenados anteriormente.
4. Cálculo de T_+ (soma das ordens de sinal positivo) e de T_- (soma das ordens de sinal negativo).

Decisão do teste: Rejeitamos H_0 se T_+ ou T_- for menor que o valor crítico $T_{\alpha(2),n}$.

NOTA: No caso da existência de empates, o procedimento é igual ao que foi referido para o teste de WMW, ou seja, atribui-se a média das ordens que as duas observações possuíam se não ocorresse empate.

Estatística Inferencial N-P

O teste de Wilcoxon (*signed-rank test*)

Exemplo: Queremos testar, a um nível de significância $\alpha=0.05$, a hipótese bilateral de não existirem diferenças entre o grau de queimadura solar com um dado protector (ombro esquerdo dos indivíduos) e sem protector solar (ombro direito dos indivíduos).

H_0 : Não há diferenças entre os graus de queimadura solar (o protector não faz efeito).

H_1 : Há diferenças entre os graus de queimadura solar (o protector faz efeito).

Indivíduo	Ombro sem protector	Ombro com protector	Diferenças	Ordem das diferenças absolutas	Ordens afectadas de sinal
1	142	138	4	4.5	4.5
2	140	136	4	4.5	4.5
3	144	147	-3	3	-3
4	144	139	5	7	7
5	142	143	-1	1	-1
6	146	141	5	7	7
7	149	143	6	9.5	9.5
8	150	145	5	7	7
9	142	136	6	9.5	9.5
10	148	146	2	2	2

$$n = 10$$

$$T_+ = 4.5+4.5+7+7+9.5+7+9.5+2 = 51$$

$$T_- = 3+1 = 4$$

$$T_{0.05(2),10} = 8 \text{ (tabela)}$$

Estatística Inferencial N-P

O teste de Wilcoxon (*signed-rank test*)

Exemplo (continuação): Consultemos então uma tabela de valores críticos da distribuição T de Wilcoxon, para um nível de significância de 0.05 e $n = 10$:

TABLE B.11 CRITICAL VALUES OF THE WILCOXON T DISTRIBUTION

n	$\alpha(2) = 0.50$	0.20	0.10	0.05	0.02	0.01
	$\alpha(1) = 0.25$	0.10	0.05	0.025	0.01	0.005
4	2	0				
5	4	2	0			
6	6	3	2	0		
7	9	5	3	2	0	
8	12	8	5	3	1	0
9	16	10	8	5	3	1
10	20	14	10	8	5	3
11	24	17	13	10	7	5
					9	7

(Continuar)

Decisão:

Uma vez que o valor crítico do teste é $U_{0.05(2)(10)} = 8$ (tabela), rejeitamos H_0 pois $T_- = 4 < 8$.

Concluímos que, a um nível de significância $\alpha=0.05$, o protector solar faz efeito.

Estatística Inferencial N-P

Análise de Dados

O teste de Wilcoxon (*signed-rank test*)

Procedimento em SPSS

Introdução dos dados

À semelhança do que já vimos para testes t a duas amostras emparelhadas, cada variável contém os valores de cada amostra.

Menu: Analyze
Nonparametric Tests
2 Related Samples Variable 1 -- Variable 2
Test Type Wilcoxon
OK

97

Lista de alguns *sites* de ajuda *online* em Estatística e Análise de dados:

- <http://davidmlane.com/hyperstat/index.html>
- <http://math.about.com/cs/statistics/>
- <http://www.meandeviation.com/tutorials/stats/>
- <http://biology.nebrwesleyan.edu/empiricist/sources/tips/stats1.html>
- <http://www.math.ist.utl.pt/stat/>

98



Diapositivos de apoio à disciplina de Análise Avançada de Dados

Docente: Pedro Sousa (pedro.sousa.mrm@gmail.com)

99

Estatística? Para quê?...

Relembremos os conceitos básicos

A Estatística é a *ferramenta* que nos permite conhecer melhor as características de um dado conjunto de indivíduos, ao qual chamamos **população**, e também tirar conclusões válidas acerca dessa população. Uma vez que a análise de todos os indivíduos da população é, na maioria das vezes, impraticável (ver exemplos abaixo), pode-se estudar apenas uma pequena parte dessa população à qual chamamos **amostra**. O objectivo é **inferir** certos factos da população, a partir dos resultados observados na amostra. Este processo denomina-se **inferência estatística**. Ao processo de obtenção de amostras chamamos **amostragem**.

Neste contexto, temos dois grandes tipos de aplicações da estatística: a **Estatística Descritiva**, quando apenas se pretende explorar e conhecer melhor um conjunto de dados, e a **Estatística Inferencial**, quando se pretende inferir, ou seja, tirar conclusões acerca da população, com base nas características da amostra.

100

Estatística? Para quê?...

Relembremos os conceitos básicos

Exemplo 1: Pretende-se tirar conclusões acerca das alturas de 20 000 alunos que frequentam uma dada Universidade (**população**), com base numa **amostra** de 100 desses alunos.

Exemplo 2: Pretende-se saber se a aplicação de determinada vacina aos habitantes de um país (**população**) é eficiente ou não, com base nos resultados da aplicação da vacina a 200 desses habitantes (**amostra**).

Deste modo, a estatística é hoje, e cada vez mais, fundamental em Medicina, Farmacologia, Ciências da comunicação, Engenharias, Biologia, Pescas, Ecologia, Psicologia e em todas as muitas outras áreas científicas.

101

Medidas de associação

Estatística:

I. Covariância e Coeficiente de Correlação Linear (Pearson)

Considerando que é sugerida a existência de uma relação linear (positiva ou negativa) entre duas variáveis métricas X e Y , torna-se necessário quantificar essa relação. Neste contexto, definamos o conceito de **covariância**. Do mesmo modo que a variância de uma dada variável consiste na soma média de quadrados dos desvios das observações relativamente à sua média, a **covariância** consiste no produto médio dos desvios das observações relativos às respectivas médias:

$$Cov(X, Y) = \sum (x_i - \bar{x})(y_i - \bar{y}) / (n - 1)$$

A covariância é uma medida do grau de relação entre duas variáveis, mas não é um índice, pois depende das unidades das variáveis. Se padronizarmos a $cov(X, Y)$ pelo produto dos desvios padrão amostrais obtemos o **coeficiente de correlação linear de Pearson** (r):

102

Covariância e Coeficiente de Correlação Linear (Pearson)

Análise Avançada de Dados

$$r = \frac{\text{cov}(X, Y)}{s_x s_y} = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum (x_i - \bar{x})^2 \sum (y_i - \bar{y})^2}}$$

$$r = \frac{n \sum x_i y_i - (\sum x_i)(\sum y_i)}{\sqrt{n \sum x_i^2 - (\sum x_i)^2} \sqrt{n \sum y_i^2 - (\sum y_i)^2}}$$

O coeficiente de Pearson varia de -1 (correlação negativa perfeita), a 1 (correlação positiva perfeita), significando um valor de 0 ausência de correlação linear.

Para verificar a significância de um dado valor de r , recorre-se à respectiva tabela de valores críticos. Em alternativa, calcula-se a estatística t com base no número de pares de observações e no valor de r . Esta estatística distribui-se como um t de *Student* com $(n-2)$ graus de liberdade, sob H_0 .

Em SPSS a significância de r é tabelada e basta-nos comparar este valor com α (regra universal).

103

Covariância e Coeficiente de Correlação Linear (Pearson)

Análise Avançada de Dados

Etapas de análise em SPSS

1. Descrever graficamente a relação entre as duas variáveis em estudo (nuvem de pontos):

Graphs
Scatter...
Simple Define **y axis** **x axis** (label cases by)
OK

2. Se é sugerida a existência de uma relação linear positiva ou negativa entre as variáveis, calcular o coeficiente de Pearson e a respectiva significância:

Analyze
Correlate
Bivariate
variables **correlation coefficient**
OK

104

Covariância e Coeficiente de Correlação Linear (Pearson)

Análise Avançada de Dados

Exemplo: amostra de 8 crianças e adolescentes dos quais se registou a idade e a estatura (cm) - DADOS:

	crianças	idade	estatura
1	1	14	155.00
2	2	6	120.00
3	3	11	145.00
4	4	7	123.00
5	5	13	161.00
6	6	12	150.00
7	7	2	91.00
8	8	4	103.00

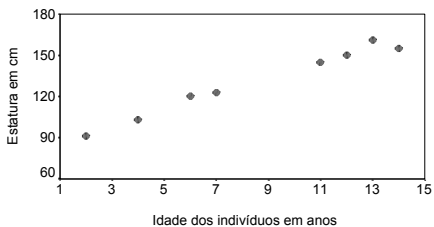
⇒ Nuvem de pontos: →

105

Covariância e Coeficiente de Correlação Linear (Pearson)

Análise Avançada de Dados

Exemplo: amostra de 8 crianças e adolescentes dos quais se registou a idade e a estatura (cm) – NUVEM DE PONTOS:



106

Covariância e Coeficiente de Correlação Linear (Pearson)

Análise Avançada de Dados

Exemplo: amostra de 8 crianças e adolescentes dos quais se registou a idade e a estatura (cm) – SPSS DIALOG BOX:



107

Covariância e Coeficiente de Correlação Linear (Pearson)

Análise Avançada de Dados

Exemplo: amostra de 8 crianças e adolescentes dos quais se registou a idade e a estatura (cm) - RESULTADOS:

Correlations

		IDADE	ALTURA
IDADE	Pearson Correlation	1	.990**
	Sig. (2-tailed)	.	.000
	N	8	8
ALTURA	Pearson Correlation	.990**	1
	Sig. (2-tailed)	.000	.
	N	8	8

** . Correlation is significant at the 0.01 level (2-tailed).

108



Correlação não-paramétrica: Coeficiente de Spearman

Se os nossos dados se afastam bastante da normalidade, o coeficiente de Pearson não é, na generalidade, aplicável. A alternativa mais conhecida não-paramétrica ao coeficiente de Pearson é o coeficiente ordinal de Spearman, o qual utiliza as ordens das observações de cada variável. O teste pode ser resumido nos seguintes procedimentos:

1. Atribuição de ordens a cada uma das variáveis, separadamente (ao contrário do efectuado para o teste não-paramétrico de Wilcoxon para duas amostras emparelhadas).
2. Calcular o vector de diferenças entre os pares de ordens.

109



Correlação não-paramétrica: Coeficiente de Spearman

3. Calcular a estatística do teste, dada por:

$$r_s = 1 - \frac{6 \sum d_i^2}{n^3 - n}$$

Sendo: " r_s " o coeficiente amostral de Spearman, " d_i^2 " o quadrado das diferenças das ordens e " n " o número de pares de observações.

4. Decisão do teste: comparação de r_s com o respectivo valor crítico tabelado (α, n). Rejeita-se H_0 se o valor calculado for superior ao tabelado. Em SPSS, utiliza-se a regra universal de decisão (*).

110



Correlação não-paramétrica: Coeficiente de Spearman

Exemplo 1:

Considere-se uma dada amostra bi-variada (X_1, X_2) $n = 7$.
Adopte-se $\alpha = 0.05$.

X_1	X_2	Ordens		Diferença ($O_{x1} - O_{x2}$)	Diferença ²
		O_{x1}	O_{x2}		
4.1	12.0	2	2	0	0
4.0	13.1	1	3	-2	4
6.0	11.0	3	1	2	4
6.3	16.2	4	4	0	0
9.8	19.7	5	5	0	0
12.1	20.8	7	6	1	1
11.3	21.4	6	7	-1	1

O somatório das diferenças elevadas ao quadrado é, neste caso, 10. Aplicando a fórmula da página anterior, $r_s = 0.821$. Uma vez que o valor crítico tabelado é **0.786** (teste bilateral), rejeita-se H_0 pois r_s é superior ao valor crítico, ou seja, as duas variáveis são correlacionadas. **Em SPSS, utiliza-se a regra universal de decisão (*)**.

111

Correlação não-paramétrica: Coeficiente de Spearman

Exemplo 2: amostra de 8 crianças e adolescentes dos quais se registou a idade e a estatura (cm) – SPSS DIALOG BOX:



112

Correlação não-paramétrica: Coeficiente de Spearman

Exemplo 2: amostra de 8 crianças e adolescentes dos quais se registou a idade e a estatura (cm) - RESULTADOS:

Correlations				
Spearman's rho	IDADE	Correlation Coefficient	1.000	.976**
		Sig. (2-tailed)	.	.000
		N	8	8
	ALTURA	Correlation Coefficient	.976**	1.000
		Sig. (2-tailed)	.000	.
		N	8	8

**. Correlation is significant at the .01 level (2-tailed).

113

Medidas de associação Estatística:

II. Análise de Tabelas de Contingência: Teste de Qui-quadrado

A secção anterior abordou métodos de quantificação do grau de associação entre duas variáveis métricas, utilizando métodos paramétricos e não-paramétricos.

No caso das duas (ou mais) variáveis para as quais se pretende estudar a relação serem de natureza categórica (nominais ou ordinais) o procedimento é, consequentemente, diferente, o qual se abordará na presente secção.

114

Análise de Tabelas de Contingência: Teste de Qui-quadrado

Considere-se, como exemplo, que se fez um questionário a 15 pessoas no qual se registou o sexo (feminino = 0, masculino = 1) e se o indivíduo fuma (1) ou não (0).

Os resultados (já numa folha de dados do SPSS) foram os apresentados ao lado (códigos e respectivos *labels*).

O **objectivo** é verificar se existe uma relação significativa entre sexo e fumador, isto é, se as mulheres fumam mais que os homens, ou vice-versa.

Códigos		Labels	
sexo	fumador	sexo	fumador
1	0	masc.	não
0	0	fem.	não
0	0	fem.	não
0	1	fem.	sim
1	0	masc.	não
1	0	masc.	não
0	1	fem.	sim
1	0	masc.	não
1	0	masc.	não
0	1	fem.	sim
0	1	fem.	sim
0	1	fem.	sim
1	1	masc.	sim
1	0	masc.	não
0	1	fem.	sim

115

Análise de Tabelas de Contingência: Teste de Qui-quadrado

O cruzamento das duas variáveis, isto é, a contagem do número de observações que correspondem a cada par de categorias, resulta numa **TABELA DE CONTINGÊNCIA**. Neste exemplo, esta tabela assume a forma:

SEXO * FUMADOR Crosstabulation

Count

		FUMADOR		Total
		não	sim	
SEXO	fem.	2	6	8
	masc.	6	1	7
Total		8	7	15

116

Análise de Tabelas de Contingência: Teste de Qui-quadrado

Neste ponto, urge as seguintes observações:

- Já estimámos a correlação entre variáveis métricas (contínuas), quer através do coeficiente linear de Pearson ou do coef. não-paramétrico de Spearman. Neste caso, as variáveis não são métricas mas sim categóricas (nominais).
- Na tabela de contingência anterior, além das contagens existem os totais marginais e o número total de observações. Os totais marginais das linhas (variável sexo) dão-nos o número total de indivíduos do sexo masculino e feminino. Os totais marginais das colunas (variável fumador) dão-nos o número total de indivíduos que fumam e que não fumam.

117

Análise de Tabelas de Contingência: Teste de Qui-quadrado

Mais observações:

- O número total de observações corresponde, assim, tanto à soma das observações no "interior" da tabela, como à soma dos totais marginais das linhas ou das colunas.
- Tratando-se de variáveis categóricas (nominais), o número de categorias pode ser 2, ou mais. Assim, atribui-se a nomenclatura **tabela de contingência 2 x 2** (como neste exemplo), **2 x 3**, **4 x 4**, etc., de acordo com o número de categorias que cada variável possui (ver exemplo seguinte).

118

Análise de Tabelas de Contingência: Teste de Qui-quadrado

Exemplo de uma tabela de contingência 2 x 4:

Sexo	Cor do cabelo				Total
	Negro	Castanho	Louro	Ruivo	
Masculino	32 (29.0)	43 (36.0)	16 (26.67)	9 (8.33)	100
Feminino	55 (58.0)	65 (72.0)	64 (53.33)	16 (16.67)	200
Total	87	108	80	25	300

Pergunta-se: "o que são os números entre parêntesis?"...
Para responder, voltemos ao primeiro exemplo:

119

Análise de Tabelas de Contingência

			FUMADOR		Total
			não	sim	
SEXO	fem.	número de observações	2	6	8
		% do total de "sexo"	25 ^{*1}	75 ^{*2}	100
		% do total de "fumador"	25	85.7 ^{*3}	53.3 ^{*7}
	masc.	número de observações	6	1	7
		% do total de "sexo"	85.7 ^{*3}	14.3 ^{*4}	100
		% do total de "fumador"	75	14.3	46.7 ^{*8}
Total		número de observações	8	7	15
		% do total de "sexo"	53.3 ^{*5}	46.7 ^{*6}	100
		% do total de "fumador"	100	100	100

Legenda:

- *1: $(2 / 8) * 100 = 25\%$ (% do número de indivíduos do **sexo fem.** que **não fumam**)
- *2: $(6 / 8) * 100 = 75\%$ (% do número de indivíduos do **sexo fem.** que **fumam**)
- *3: $(6 / 7) * 100 = 85.7\%$ (% do número de indivíduos do **sexo masc.** que **não fumam**)
- *4: $(1 / 7) * 100 = 14.3\%$ (% do número de indivíduos do **sexo masc.** que **fumam**)
- *5: $(8 / 15) * 100 = 53.3\%$ (% do número **total** de indivíduos que **não fumam**)
- *6: $(7 / 15) * 100 = 46.7\%$ (% do número **total** de indivíduos que **fumam**)
- *7: $(8 / 15) * 100 = 53.3\%$ (% do número **total** de indivíduos do **sexo feminino**)
- *8: $(7 / 15) * 100 = 46.7\%$ (% do número **total** de indivíduos do **sexo masculino**)

120

Análise de Tabelas de Contingência: Teste de Qui-quadrado

O princípio de análise tem a ver com um raciocínio do tipo:

Havendo um total de 53.3% de indivíduos do sexo feminino, **ESPERA-SE** que, se não existir relação entre as duas variáveis, 53.3% do número total de **não fumadores** serão mulheres (53.3% de 8 = **4.3**) e 46.7 % serão homens (46.7% de 8 = **3.7**). Do mesmo modo, 53.3% do número total de **fumadores** serão mulheres (53.3% de 7 = **3.7**) e 46.7 % serão homens (46.7% de 7 = **3.3**).

Estes são os **VALORES ESPERADOS** do cruzamento das categorias das duas variáveis, atrás assinalados entre parêntesis. Neste exemplo, assinalem-se na tabela seguinte.

O teste estatístico que permite verificar se existe uma relação significativa entre as duas variáveis compara os valores observados com os esperados (ver formulação seguinte).

121

Análise de Tabelas de Contingência: Teste de Qui-quadrado

Tabela de contingência com valores observados (*count*) e esperados (*expected count*):

SEXO * FUMADOR Crosstabulation

			FUMADOR		Total
			não	sim	
SEXO	fem.	Count	2	6	8
		Expected Count	4.3	3.7	8.0
	masc.	Count	6	1	7
		Expected Count	3.7	3.3	7.0
Total	Count		8	7	15
	Expected Count		8.0	7.0	15.0

122

Análise de Tabelas de Contingência: Teste de Qui-quadrado

O valor do teste (Qui-quadrado de Pearson) é-nos dado por:

$$\chi^2 = \sum_{i=1}^I \sum_{j=1}^c \frac{(O_{ij} - E_{ij})^2}{E_{ij}}$$

Onde O_{ij} são os valores observados e E_{ij} os esperados.

No caso do nosso exemplo este valor seria obtido do seguinte modo:

$$\chi^2 = \frac{(2-4.3)^2}{4.3} + \frac{(6-3.7)^2}{3.7} + \frac{(6-3.7)^2}{3.7} + \frac{(1-3.3)^2}{3.3} \approx 5.6$$

123

Análise de Tabelas de Contingência: Teste de Qui-quadrado

O grau de associação entre as variáveis é-nos dado pelo coeficiente de associação (ou de contingência), calculado com base no valor de qui-quadrado anterior:

$$C = \sqrt{\frac{\chi^2}{n + \chi^2}}$$

Este valor, situa-se entre 0 e 1 (máxima associação).

A **significância** do valor de qui-quadrado (*chi-square*) e do coeficiente de associação (*contingency coefficient*) é-nos dado pelo SPSS. A regra de decisão estatística* é posteriormente aplicada. As hipóteses previamente estabelecidas são:

H₀: As variáveis são independentes (não há relação entre as mesmas).

H₁: Há relação significativa entre as variáveis.

124

Análise de Tabelas de Contingência: Teste de Qui-quadrado

Procedimento em SPSS

A introdução dos dados:

Há duas situações distintas: ou os dados são introduzidos *em bruto*, isto é, segundo o exemplo do *slide* 114, ou num formato mais parecido com a tabela de contingência, que já resume a informação anterior. Vejamos esta última alternativa, usando os mesmo dados do *slide* 114:

Códigos

Labels

sexo	fumador	observ
0	0	2
0	1	6
1	0	6
1	1	1

Ou

sexo	fumador	observ
fem.	não fuma	2
fem.	fuma	6
masc.	não fuma	6
masc.	fuma	1

125

Análise de Tabelas de Contingência: Teste de Qui-quadrado

Procedimento em SPSS

A introdução dos dados: dados ponderados

Neste caso os dados estão ponderados pela variável "observ", isto é, lendo as linhas da tabela, o número de indivíduos do sexo feminino (0) que não fumam (0) é 2; o número de indivíduos do sexo feminino (0) que fumam (1) é 6; etc.

Por este motivo, neste caso, há que "informar" o SPSS que existe esta ponderação:

Data

Weight cases

Weight cases by nome da variável OK

Deste modo, este procedimento não precisa de ser efectuado se se utilizar uma tabela de dados do primeiro tipo (*slide* 114).

126

Análise de Tabelas de Contingência: Teste de Qui-quadrado

Exemplo em SPSS
Análise de qui-quadrado

Procedamos agora à análise de contingência:

Analyze

Descriptive statistics

Crosstabs

Row(s) **nome da variável linha**

Column(s) **nome da variável coluna**

Statistics **chi-square contingency coefficient** CONTINUE

Cells **observed expected** CONTINUE

OK

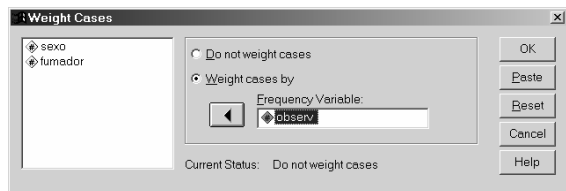
127

Análise de Tabelas de Contingência: Teste de Qui-quadrado

Procedimento em SPSS

Ponderação de observações (*weight cases*):

Menu: Data \ Weight cases



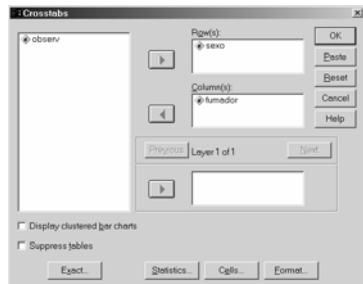
128

Análise de Tabelas de Contingência: Teste de Qui-quadrado

Procedimento em SPSS

Análise de contingência:

Menu: Analyze \ Descriptive Statistics \ Crosstabs



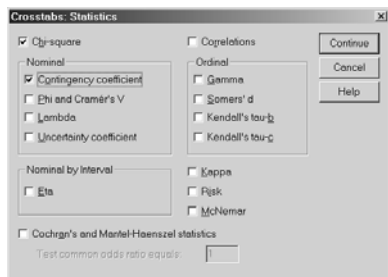
129

Análise de Tabelas de Contingência: Teste de Qui-quadrado

Procedimento em SPSS

Análise de contingência:

Menu: Analyze \ Descriptive Statistics \ Crosstabs \ Statistics



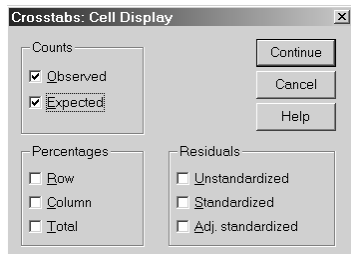
130

Análise de Tabelas de Contingência: Teste de Qui-quadrado

Procedimento em SPSS

Análise de contingência:

Menu: Analyze \ Descriptive Statistics \ Crosstabs \ Cells



131

Análise de Tabelas de Contingência: Teste de Qui-quadrado

Procedimento em SPSS

Análise de contingência:

Resultados: tabela de contingência: slide 79; tabela de valores críticos e significâncias:

Chi-Square Tests

	Value	df	Asymp. Sig. (2-sided)	Exact Sig. (2-sided)	Exact Sig. (1-sided)
Pearson Chi-Square	5.529 ^a	1	.019		
Continuity Correction ^a	3.359	1	.067		
Likelihood Ratio	5.989	1	.014		
Fisher's Exact Test				.041	.032
Linear-by-Linear Association	5.161	1	.023		
N of Valid Cases	15				

132

Análise de Tabelas de Contingência: Teste de Qui-quadrado

Procedimento em SPSS

Análise de contingência:

Resultados: tabela de valores críticos e significâncias (discussão).

Estão realçados o valor de qui-quadrado (5.529) e a respectiva significância (0.019): para um alfa de 0.05, rejeitamos H_0 , isto é, as variáveis são correlacionadas.

Observando a tabela de contingência verifica-se que o valor observado de mulheres que fumam (6) é superior ao esperado (3.7) e que o valor observado de homens que fumam (1) é inferior ao esperado (3.3).

O teste concluiu que estas diferenças são significativas ($\alpha=0.05$), ou seja, pode-se afirmar que, a um nível de significância $\alpha=0.05$, as mulheres fumam mais que os homens, na população amostrada.

133

Análise de Tabelas de Contingência: Teste de Qui-quadrado

Procedimento em SPSS

Análise de contingência:

Resultados: tabela do coeficiente de contingência:

Symmetric Measures

	Value	Approx. Sig.
Nominal by Nominal Contingency Coefficient	.519	.019
N of Valid Cases	15	

Estão realçados o valor do coeficiente de contingência (0.519) e a respectiva significância (0.019): para um alfa de 0.05, este valor de C é significativo.

134

ANÁLISE DE VARIÂNCIA

(ANOVA - Analysis Of Variance)

Comparações entre $k > 2$ amostras

INTRODUÇÃO

Verificámos anteriormente que, quando queremos comparar as médias de duas amostras referentes a uma dada variável, utilizamos o teste t . No entanto, frequentemente queremos comparar mais que duas amostras. Poderá ser tentador utilizar testes t a todos os possíveis pares de amostras, mas este procedimento para $k > 2$ amostras é inválido: o que sucede é que, enquanto entre duas amostras a probabilidade de cometermos um erro de tipo I (rejeitar H_0 quando verdadeira) é α (por exemplo 5%), se aplicarmos este procedimento a mais que duas amostras, é demonstrado (Zar, 1984), que este erro aumenta proporcionalmente em relação ao número de amostras que consideramos. Por exemplo, se estabelecemos $\alpha=5\%$ e são comparadas 3 médias, temos 13% de probabilidade de concluir erradamente uma diferença entre as duas médias mais extremas (erro tipo I). Para 10 médias esta probabilidade aumenta para 63% e considerando 20 médias temos 93% de probabilidade de cometer um erro tipo I.

135

ANÁLISE DE VARIÂNCIA

(ANOVA - Analysis Of Variance)

Comparações entre $k > 2$ amostras**INTRODUÇÃO**

Neste contexto, a Análise de Variância (ANOVA) permite verificar qual é o efeito de uma ou mais variáveis independentes (factores) numa variável dependente. A questão central da ANOVA consiste em saber se as populações têm ou não, estatisticamente, médias iguais. Formulamos então as hipóteses nula e alternativa:

$$H_0: \mu_1 = \mu_2 = \dots = \mu_k$$

H_1 : As médias das populações não são iguais.

Um aspecto importante da hipótese alternativa é que a mesma contempla não só o caso em que existe desigualdade estatística entre todas as k médias, mas também casos em que apenas uma média difira do conjunto das restantes $k-1$.

136

ANÁLISE DE VARIÂNCIA

(ANOVA - Analysis Of Variance)

Comparações entre $k > 2$ amostras**INTRODUÇÃO**

Na análise de variância os factores podem ser fixos (determinados *a priori*) ou aleatórios (escolhidos aleatoriamente). Consoante a aleatoriedade ou não dos factores temos 3 tipos de ANOVA: Tipo I (factores fixos), Tipo II (factores aleatórios) e Tipo III (factores fixos e aleatórios). Considerando o número de factores, temos ANOVA a 1 factor (*one-way ANOVA*) a 2 factores (*Two-way ANOVA*), etc.

A ideia do teste estatístico (F) consiste em comparar a variação entre os k grupos (amostras) com a variação dentro dos grupos. Se a variação entre os grupos for grande comparativamente à variação dentro dos grupos, podemos dizer que há diferenças entre as médias que queremos comparar.

137

ANÁLISE DE VARIÂNCIA

(ANOVA - Analysis Of Variance)

Comparações entre $k > 2$ amostras**INTRODUÇÃO**

Mais concretamente, considerando uma ANOVA a 1 factor fixo (Tipo I), o teste estatístico consiste no cálculo do seguinte valor de F :

$$F_{(k-1, n-k)} = \frac{\sum_{i=1}^k n_i (\bar{x}_i - \bar{x})^2 / k - 1}{\sum_{i=1}^k \sum_{j=1}^{n_i} (x_{ij} - \bar{x}_i)^2 / n - k}$$

k = número de grupos
 n = número total de observações
 n_i = número de observações do grupo i
 x_{ij} = observação j do grupo i
 $\bar{x}_i$ = média das observações do grupo i
 $\bar{x}$ = média de todas as n observações

138

ANÁLISE DE VARIÂNCIA(ANOVA - *Analysis Of Variance*)**Comparações entre $k > 2$ amostras****INTRODUÇÃO**

É frequente elaborar-se o cálculo da estatística F atrás referida, por etapas, num **quadro de ANOVA**. Para uma ANOVA a um factor fixo, o quadro genérico é representado abaixo.

Fonte de variação	Soma de quadrados	Graus de liberdade	Quadrados Médios	F
Entre grupos	SQE	$k - 1$	$QME = SQE / (k - 1)$	QME / QMR
Dentro grupos (resíduo, erro)	SQR	$n - k$	$QMR = SQR / (n - k)$	
Total	SQT	$n - 1$		

Como pode ser observado na fórmula da página anterior, SQE e SQR são dados pelas expressões abaixo:

$$SQE = \sum_{i=1}^k n_i (\bar{x}_i - \bar{x})^2 \quad SQR = \sum_{i=1}^k \left[\sum_{j=1}^{n_i} (x_{ij} - \bar{x}_i)^2 \right]$$

139

**Comparações múltiplas
posteriores à ANOVA:
O teste de Tukey**

Supondo que o valor da estatística F é significativo (α), podemos querer saber entre que grupos (amostras) é que existem diferenças. Isto porque, basta duas de entre as k amostras serem diferentes estatisticamente (embora todas as outras iguais) para rejeitarmos H_0 , isto é, para o valor de F calculado ser significativo. O teste mais conhecido para verificar estas diferenças é o de **TUKEY**:

Consideremos que, para k grupos, ocorrem $k(k-1)/2$ combinações de pares de grupos. Por exemplo, para 3 grupos (A, B e C) temos $3(3-1)/2 = 3$ pares de grupos, nomeadamente, A-B, A-C e B-C. O procedimento do teste de Tukey é representado na tabela abaixo:

140

Comparação	Diferenças	SE	q	$q(\alpha, n, k)$	Decisão do teste
A vs. B	$\bar{X}_{A,B} = \bar{X}_A - \bar{X}_B$	$SE_{A,B} = \sqrt{QMR / n_{A,B}}$	$q_{A,B} = \bar{X}_{A,B} / SE_{A,B}$	Tabela (q)	Rejeito H_0 se $q > q_c$
A vs. C	$\bar{X}_{A,C} = \bar{X}_A - \bar{X}_C$	$SE_{A,C} = \sqrt{QMR / n_{A,C}}$	$q_{A,C} = \bar{X}_{A,C} / SE_{A,C}$	Tabela (q)	Rejeito H_0 se $q > q_c$
B vs. C	$\bar{X}_{B,C} = \bar{X}_B - \bar{X}_C$	$SE_{B,C} = \sqrt{QMR / n_{B,C}}$	$q_{B,C} = \bar{X}_{B,C} / SE_{B,C}$	Tabela (q)	Rejeito H_0 se $q > q_c$

**Comparações múltiplas
posteriores à ANOVA:
O teste de Tukey**

No quadro anterior, $n_{A,B}$ é o número de observações comuns aos grupos A e B. No caso deste número ser diferente, utiliza-se uma ponderação [ver Zar (1984)]. QMR é o valor do quadrado médio dos resíduos (variação dentro dos grupos) da ANOVA. Considerando 3 grupos (A, B, C) os resultados possíveis seriam os seguintes:

$\mu_A \neq \mu_B \neq \mu_C$	Todos os grupos diferentes
$\mu_A = \mu_B \neq \mu_C$	Grupos A e B sem $\neq$ entre si mas $\neq$ de C
$\mu_A \neq \mu_B = \mu_C$	Grupos B e C sem $\neq$ entre si mas $\neq$ de A
$\mu_A = \mu_C \neq \mu_B$	Grupos A e C sem $\neq$ entre si mas $\neq$ de B
$\mu_A = \mu_C = \mu_B$	Todos os grupos sem $\neq$ significativas

Em SPSS, os resultados do teste de Tukey são apresentados numa tabela relacionando cada par de grupos, tal como será analisado posteriormente.

141

ANÁLISE DE VARIÂNCIA

(ANOVA - Analysis Of Variance)

Análise Avançada de Dados

Pressupostos da ANOVA

Tal como para o teste t para duas amostras, temos que assumir que as observações de cada amostra se distribuem normalmente, e que possuem variâncias iguais (homoscedasticidade). Se ocorrerem desvios muito grandes a estes pressupostos, podemos transformar os dados, ou utilizar testes não-paramétricos, que, como já vimos anteriormente, não dependem deste tipo de pressupostos.

142

ANÁLISE DE VARIÂNCIA

(ANOVA - Analysis Of Variance)

Análise Avançada de Dados

Procedimentos em SPSS

1. Introdução dos dados:

Tal como para testes de comparação entre duas amostras independentes, uma variável contém os valores das amostras (grupos) e a outra a indicação 1... k correspondente aos k grupos.

2. Testes à Normalidade das observações em cada grupo

Analyze
Descriptive Statistics
Explore
Dependent list (variável com os valores)
Factor list (variável com a indicação dos grupos)
Plots **Normality plots with tests** Continue OK

143

ANÁLISE DE VARIÂNCIA

(ANOVA - Analysis Of Variance)

Análise Avançada de Dados

Procedimentos em SPSS

3. Anova a 1 factor e teste de Tukey de comparações posteriores (Post-Hoc):

Analyze
Compare Means
One-Way Anova
Dependent list (variável com os valores)
Factor list (variável com a indicação dos grupos)
Post-Hoc
Equal variances assumed **Tukey**
Significance level (valor de alfa) Continue
Options
Homogeneity of variance
Means plot
Continue OK

144

ANÁLISE DE VARIÂNCIA

(ANOVA - Analysis Of Variance)

Exemplo em SPSS

1. Dados:

Grupo I	Grupo II	Grupo III
1.53	3.15	8.18
1.61	3.96	5.64
3.75	3.59	7.36
2.89	1.89	8.82
	4.32	



Em SPSS:

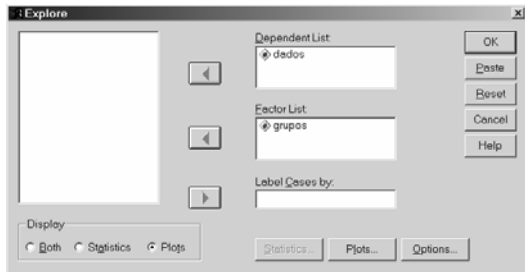
	dados	grupos
1	1.53	1
2	1.61	1
3	3.75	1
4	2.89	1
5	3.15	2
6	3.96	2
7	3.59	2
8	1.89	2
9	4.32	2
10	8.18	3
11	5.64	3
12	7.36	3
13	8.82	3

145

ANÁLISE DE VARIÂNCIA

(ANOVA - Analysis Of Variance)

Exemplo em SPSS

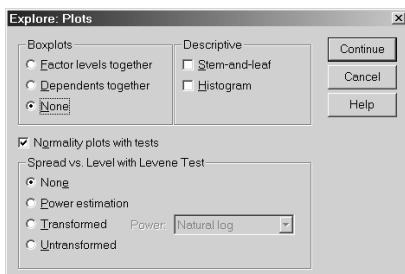
2. Testes à Normalidade : sub-menu *Explore*

146

ANÁLISE DE VARIÂNCIA

(ANOVA - Analysis Of Variance)

Exemplo em SPSS

2. Testes à Normalidade: sub-menu *Explore / Plots*

147

ANÁLISE DE VARIÂNCIA

(ANOVA - Analysis Of Variance)

Exemplo em SPSS

2. Testes à Normalidade: Resultados

Tests of Normality						
GRUPOS	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
DADOS 1	.282	4		.880	4	.338
2	.203	5	.200 [*]	.927	5	.575
3	.209	4		.948	4	.707

* This is a lower bound of the true significance.

a. Lilliefors Significance Correction

Relembrando quais são as hipóteses...

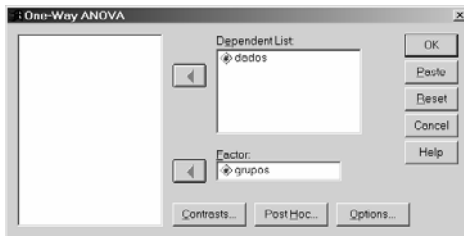
Decidimos que **HÁ NORMALIDADE** em todos os grupos
(regra univ. de decisão estatística em SPSS...)

148

ANÁLISE DE VARIÂNCIA

(ANOVA - Analysis Of Variance)

Exemplo em SPSS

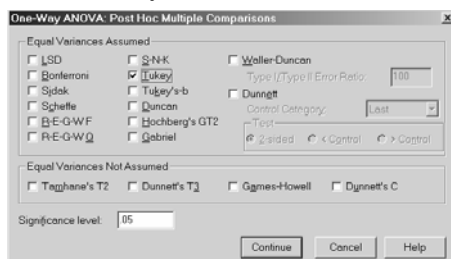
3. ANOVA a 1 factor e teste de Tukey: sub-menu One-way ANOVA

149

ANÁLISE DE VARIÂNCIA

(ANOVA - Analysis Of Variance)

Exemplo em SPSS

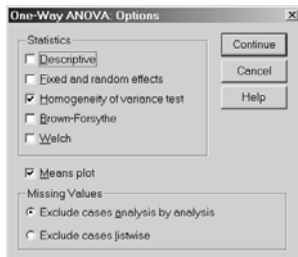
3. ANOVA a 1 factor e teste de Tukey:
sub-menu *One-way ANOVA / Post Hoc*

150

ANÁLISE DE VARIÂNCIA

(ANOVA - Analysis Of Variance)

Exemplo em SPSS

3. ANOVA a 1 factor e teste de Tukey:sub-menus *One-way ANOVA / Options*

151

ANÁLISE DE VARIÂNCIA

(ANOVA - Analysis Of Variance)

Exemplo em SPSS

3. ANOVA a 1 factor e teste de Tukey - RESULTADOS:

Homogeneidade de variâncias e quadro da ANOVA

Test of Homogeneity of Variances

DADOS				
Levene Statistic	df1	df2	Sig.	
.343	2	10	.718	

ANOVA

DADOS

	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	58.890	2	29.445	23.264	.000
Within Groups	12.657	10	1.266		
Total	71.546	12			

152

ANÁLISE DE VARIÂNCIA

(ANOVA - Analysis Of Variance)

Exemplo em SPSS

3. ANOVA a 1 factor e teste de Tukey - RESULTADOS:

Quadro de Tukey

Multiple Comparisons

Dependent Variable: DADOS

Tukey HSD

(I) GRUPOS	(J) GRUPOS	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
1	2	-.9370	.75468	.457	-3.0058	1.1318
	3	-5.0550*	.79551	.000	-7.2357	-2.8743
2	1	.9370	.75468	.457	-1.1318	3.0058
	3	-4.1180*	.75468	.001	-6.1868	-2.0492
3	1	5.0550*	.79551	.000	2.8743	7.2357
	2	4.1180*	.75468	.001	2.0492	6.1868

*. The mean difference is significant at the .05 level.

153

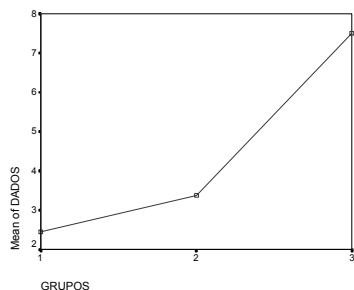
ANÁLISE DE VARIÂNCIA

(ANOVA - Analysis Of Variance)

Exemplo em SPSS

**3. ANOVA a 1 factor
e teste de Tukey –
RESULTADOS:**

Quadro de média
dos grupos
(*means plot*)



154

ANÁLISE DE VARIÂNCIA

(ANOVA - Analysis Of Variance)

Exemplo em SPSS

3. ANOVA a 1 factor e teste de Tukey**CONCLUSÕES:**

1. Há homogeneidade ($\alpha = 0.05$) entre as variâncias dos grupos (quadro *homogeneity of variances* - *Sig.* = 0.718);
2. Há diferenças significativas ($\alpha = 0.05$) entre grupos (quadro da ANOVA – *Sig.* < 0.000);
3. Os grupos 1 e 2 não diferem entre si ($\alpha = 0.05$) mas diferem do grupo 3 (quadro de Tukey);
4. O grupo 3 difere por excesso (valores maiores) comparativamente aos grupos 1 e 2 (*means plot*).

155

Dados NÃO NORMAIS:**Quais os passos a seguir?**

Como já foi referido, um dos pressupostos na aplicação de testes paramétricos tais como a ANOVA, é que as observações se distribuam segundo uma lei Normal.

Quando a Análise de Variância é inviabilizada pela quebra deste pressuposto temos duas alternativas:

1. Proceder a transformações nos dados com o objectivo de atingir a Normalidade;
2. Proceder a uma Análise de Variância Não-Paramétrica: teste de Kruskal-Wallis.

156

Dados NÃO NORMAIS: Transformação

Exemplo em SPSS: dados

	Grupo1	Grupo2	Grupo3
Consideremos o conjunto de dados ao lado, referente a uma experiência com 3 diferentes níveis de um dado factor.	9.58	5.26	1.08
	9.97	7.46	1.88
	7.10	6.30	9.78
	15.96	7.54	45.00
	7.46	12.43	7.10
Analiseemos se os individuos se distribuem Normalmente em cada grupo ($\alpha = 0.05$).	7.61	4.53	9.97
	10.70	7.24	23.40
	7.69	3.63	6.62
	8.08	10.38	1.30
	7.32	12.68	
	6.23	9.78	
	5.37		

157

Dados NÃO NORMAIS: Transformação

Exemplo em SPSS: Normalidade?

Tests of Normality							
GRUPOS		Kolmogorov-Smirnov			Shapiro-Wilk		
		Statistic	df	Sig.	Statistic	df	Sig.
DADOS	1	.239	12	.057	.828	12	.020
	2	.187	11	.200	.944	11	.563
	3	.329	9	.006	.754	9	.006

Podemos observar que não se verifica a Normalidade nos grupos 1 e 3, embora se verifique no grupo 2 ($\alpha = 0.05$).

158

Dados NÃO NORMAIS: Transformação

Exemplo em SPSS: transformação – Ln (dados)

Transformemos então os dados: uma das transformações que mais se utiliza com vista a obter a Normalidade é a **logaritimização em base natural**.

Transform
Compute
Target variable (*nome da nova var. logaritimizada*)
Numeric expression : LN(*nome da var. original*)
OK

Façamos novamente o teste de Shapiro-Wilk à Normalidade, agora dos dados logaritimizados:

159

Dados NÃO NORMAIS: Transformação

Exemplo em SPSS: Normalidade dos dados transformados?

Tests of Normality						
	GRUPOS	Kolmogorov-Smirnov			Shapiro-Wilk	
		Statistic	df	Sig.	Statistic	df
LN_DADOS	1	.196	12	.200	.930	12
	2	.120	11	.200	.958	11
	3	.189	9	.200	.937	9

Pode verificar-se que a transformação (logaritmização) conferiu Normalidade à distribuição das observações, uma vez que não se rejeita H_0 para nenhum dos grupos, considerando $\alpha=0.05$. Havendo Normalidade, podemos proceder à Análise de Variância que nos propusemos.

160

Dados NÃO NORMAIS: Transformação

Advertências sobre a transformação de dados

1. A logaritmização não é o único processo de transformação que se pode experimentar visando atingir a Normalidade. Na realidade, a escolha do tipo de transformação a efectuar depende fortemente do tipo de dados que se possui e dos objectivos da análise (Zar, 1984).
2. Após a transformação dos dados, todos os parâmetros estimados têm de sofrer a transformação contrária à efectuada de forma a ser possível compará-los nas unidades originais.
3. Nem sempre a transformação resolve o problema. Em muitos casos a não-normalidade mantém-se e a solução passa pela utilização de testes não-paramétricos.

161

Dados NÃO NORMAIS: ANOVA Não-Paramétrica

O teste de Kruskal-Wallis ($k > 2$ amostras independentes)

Se um determinado conjunto de dados consiste em mais que duas amostras independentes (grupos) sobre um dado factor, o teste não-paramétrico de Kruskal-Wallis (K-W) pode ser aplicado para testar as diferenças entre as médias dos k grupos. O teste de K-W pode ser aplicado em todas as situações em que é possível utilizar a ANOVA paramétrica a 1 factor, e também em situações em que a ANOVA paramétrica não é aplicável tais como quando as k amostras são provenientes de populações não-Normais, ou com variâncias não homogêneas (heteroscedasticidade). O teste de K-W pode ser considerado uma extensão do teste de Mann-Whitney (M-W, já referido) para $k > 2$ amostras. Para $k = 2$ amostras, o teste de M-W é o indicado.

162

Dados NÃO NORMAIS: ANOVA Não-Paramétrica

O teste de Kruskal-Wallis ($k > 2$ amostras independentes)

O teste de K-W pode ser descrito nas seguintes etapas:

1. Atribuição das ordens das observações em cada grupo, relativamente ao total das n observações.
2. Cálculo da estatística H (fórmula abaixo) na qual R_i consiste na soma das ordens do grupo i , n = número total de observações e n_i = número de observações do grupo i .
3. **Decisão do teste:** Para pequenas amostras ($n < 15$) em que $k \leq 5$ os valores críticos do teste são dados por uma tabela própria. Para amostras maiores e/ou $k > 5$, H pode ser considerada aproximada por um χ^2 com $k - 1$ graus de liberdade.

Em SPSS utiliza-se a regra generalizada de decisão*

$$H = \frac{12}{n(n+1)} \sum_{i=1}^k \frac{R_i^2}{n_i} - 3(n+1)$$

163

Dados NÃO NORMAIS: ANOVA Não-Paramétrica

O teste de Kruskal-Wallis : procedimento em SPSS

Introdução dos dados

Tal como para a comparação paramétrica entre k amostras independentes (ANOVA) uma variável contém os valores das amostras (grupos) e a outra a indicação 1... k correspondente aos k grupos.

Estatística de K-W e respectiva significância

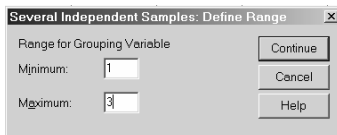
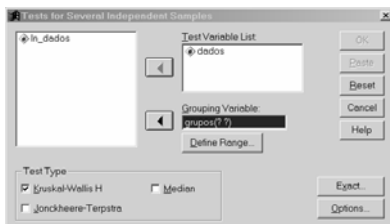
Analyze
Nonparametric Tests
k independent samples
test variable list (variável com os valores)
grouping variable (var. com a indicação dos grupos)
define range 1 k
test type **Kruskal-Wallis H**

OK

164

Dados NÃO NORMAIS: ANOVA Não-Paramétrica

O teste de Kruskal-Wallis :
procedimento em
SPSS - MENUS



165

Dados NÃO NORMAIS: ANOVA Não-Paramétrica

O teste de Kruskal-Wallis :

Resultados, utilizando o conjunto de dados atrás referido no tema *transformação de dados*.

Ranks			
	GRUPOS	N	Mean Rank
DADOS	1	12	17.96
	2	11	15.91
	3	9	15.28
	Total	32	

Neste caso não se rejeita H_0 visto que a significância do teste (0.784) é superior a α (e.g., 0.05).

Assim, não há diferenças entre os grupos.

Test Statistics ^{a,b}	
	DADOS
Chi-Square	487
df	2
Asymp. Sig.	.784

a. Kruskal-Wallis Test

b. Grouping Variable: GRUPOS

166

Comparação N-P com $k > 2$ amostras emparelhadas: o teste de Friedman

Introdução

Vimos anteriormente como, através do teste não-paramétrico de Wilcoxon, se podem comparar duas amostras emparelhadas. O teste de Friedman consiste numa generalização do teste de Wilcoxon para $k > 2$ amostras emparelhadas. Mais, tal como o teste de Kruskal-Wallis é uma alternativa não-paramétrica à ANOVA a 1 factor, o teste de Friedman consiste numa alternativa não-paramétrica para a ANOVA a 2 factores.

Deste modo, temos dois factores na experiência; um factor **a** que consiste na variável de teste, e um factor **b** que define o emparelhamento das amostras, ou seja, os indivíduos, grupos ou blocos. Consideremos o exemplo da tabela abaixo, no qual se consideram as classificações finais de 3 estudantes seleccionados aleatoriamente ($b = 3$), a 3 disciplinas diferentes ($a = 3$). Pretende-se saber se o desempenho da população de alunos da qual esta amostra foi extraída é igual em todas as disciplinas (H_0), ou se é diferente, pelo menos, em uma delas (H_1).

167

Comparação N-P com $k > 2$ amostras emparelhadas: o teste de Friedman

Introdução

Estudante	Disciplina		
	I. Aplicada 2	Inglês	Economia
A	18	8	12
B	12	10	15
C	15	14	13

Note-se que as classificações obtidas às disciplinas (variáveis - factor **a**), estão emparelhadas pelo factor **b** (estudante), por exemplo, as classificações "18", "8" e "12" estão relacionadas (emparelhadas) entre si uma vez que foram obtidas pelo mesmo estudante.

168



Comparação N-P com

$k > 2$ amostras emparelhadas: o teste de Friedman

Introdução

O teste de Friedman pode ser descrito nas seguintes etapas:

1. Atribuição das ordens (crescentes) das observações em cada bloco, ou seja, em cada linha da tabela, relativamente ao total das a observações nesse bloco (tabela de ordens das observações).

Tabela de ordens das observações

Estudante	Disciplina		
	I. Aplicada 2	Inglês	Economia
A	3	1	2
B	2	1	3
C	3	2	1
R_i	8	4	6



Comparação N-P com

$k > 2$ amostras emparelhadas: o teste de Friedman

Introdução

2. Cálculo da estatística de Friedman χ_r^2 (fórmula abaixo) na qual R_i consiste na soma das ordens do nível i da variável de teste, a = número de níveis da variável de teste e b = número de indivíduos.

$$\chi_r^2 = \frac{12}{ba(a+1)} \sum_{i=1}^a R_i^2 - 3b(a+1)$$

3. **Decisão do teste:** Para pequenas amostras ($a \leq 5$ e $b \leq 8$) os valores críticos do teste são dados por uma tabela própria de Friedman. Para amostras maiores a estatística de Friedman pode ser considerada aproximada por um χ^2 com $(a - 1)$ graus de liberdade. Rejeita-se H_0 se o valor calculado da estatística de Friedman superar o tabelado (α).

EM SPSS, UTILIZA-SE A REGRA UNIVERSAL DE DECISÃO*



Comparação N-P com

$k > 2$ amostras emparelhadas: o teste de Friedman

Procedimento em SPSS

Introdução dos dados

Tal como para a comparação não-paramétrica entre 2 amostras emparelhadas (teste de Wilcoxon) cada amostra é introduzida numa variável separada.

Estatística de Friedman e respectiva significância

Analyze
Nonparametric Tests
k related samples
test variables vars. correspondentes às amostras
test type Friedman

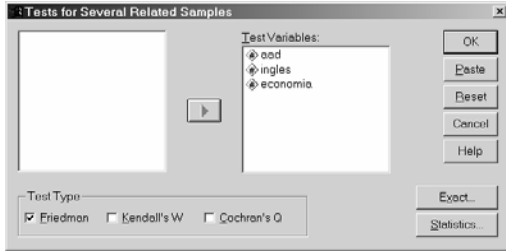
OK

Comparação N-P com

$k > 2$ amostras emparelhadas: o teste de Friedman

Exemplo em SPSS:

Menu "Analyze \ Nonparametric tests \ k related samples"



172

Comparação N-P com

$k > 2$ amostras emparelhadas: o teste de Friedman

Exemplo em SPSS:

RESULTADOS

Neste caso não se rejeita H_0 visto que a significância do teste (Asymp. Sig. = 0.264) é superior a α (e.g., 0.05).

Assim, conclui-se ($\alpha = 0.05$) que não há diferenças entre os grupos.

Nota: NESTE EXEMPLO, a pequena dimensão da amostra ($n = 3$) confere elevado grau de incerteza na rejeição de H_0 . Por este motivo, apesar de visualmente se verificar que as classificações a AAD são relativamente superiores, estatisticamente não se pode fazer esta afirmação (= não se rejeita H_0).

Ranks

	Mean Rank
AAD	2.67
INGLES	1.33
ECONOMIA	2.00

Test Statistics^a

N	3
Chi-Square	2.667
df	2
Asymp. Sig.	.264

a. Friedman Test

173

Testes baseados na Distribuição Binomial

0. Breve referência à distribuição Binomial

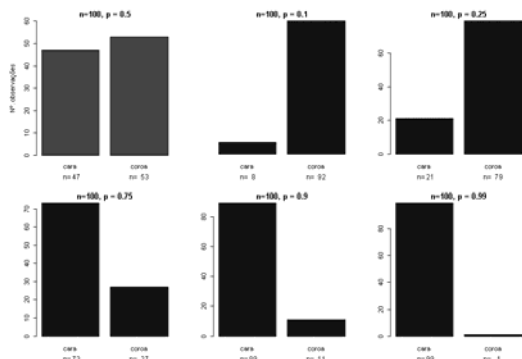
A distribuição Binomial descreve a probabilidade de se obter exactamente k sucessos, com probabilidade de sucesso p , e de insucesso $(1-p)$, numa série de eventos independentes n .

Vejamos o exemplo seguinte, no qual seis moedas são atiradas ao ar, uma delas não viciada, ou seja, com probabilidade de sucesso p (cara) = probabilidade de insucesso $q=(1-p)$ (coroa) = 0.5. As restantes moedas são viciadas de 5 modos diferentes...



174

Testes baseados na Distribuição Binomial



175

Testes baseados na Distribuição Binomial

I. O teste Binomial

Assim, se os dados consistem de n eventos independentes, cada resultado será ou "classe 1" ou "classe 2", mas nunca ambos. O número de observações nas classe 1 e 2 serão "O1" e "O2", respectivamente.



Neste ponto, cabe relembrar que os dados poderão ser agrupados ou não, sendo que, no primeiro caso haverá que efectuar uma ponderação prévia (**weight cases**) a todas as análises a efectuar.

Vejamos o exemplo seguinte:

176

Testes baseados na Distribuição Binomial

I. O teste Binomial: exemplo

Os dados ao lado correspondem a uma amostra de $n=16$ indivíduos adultos:

"Classe 1" = 0 = "s/ curso superior"
 "Classe 2" = 1 = "c/ curso superior".

Questão:

A um nível de sig. de 0.01 será 0.4 a proporção de indivíduos sem curso superior?

$$H_0: p = 0.4$$

$$H_1: p \neq 0.4$$

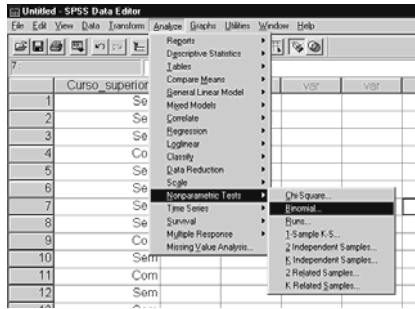
(teste bilateral)

	Curso_superior
1	Sem
2	Sem
3	Sem
4	Com
5	Sem
6	Sem
7	Sem
8	Sem
9	Com
10	Sem
11	Com
12	Sem
13	Sem
14	Sem
15	Com
16	Sem

177

Testes baseados na Distribuição Binomial

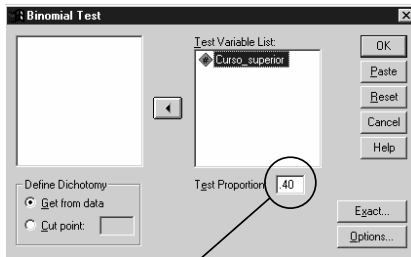
I. O teste Binomial: exemplo – procedimento em SPSS: (1)



178

Testes baseados na Distribuição Binomial

I. O teste Binomial: exemplo – procedimento em SPSS: (2)



Proporção a testar ("classe 1")

179

Testes baseados na Distribuição Binomial

I. O teste Binomial: exemplo – procedimento em SPSS: Resultado

		Category	N	Observed Prop.	Test Prop.	Exact Sig. (1-tailed)
curso_sup	Group 1	sem curso	12	,75	,4	,005
	Group 2	com curso	4	,25		
	Total		16	1,0		

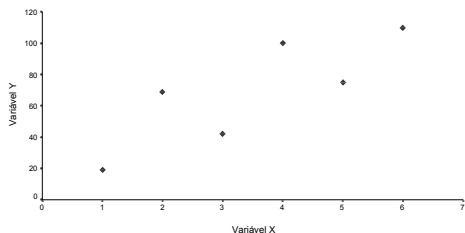
Uma vez que a sig. é menor que α (0.01), rejeita-se H_0 , ou seja, a proporção de indivíduos sem curso superior é diferente de 0.4. Pela tabela verifica-se que esta diferença é por excesso (porque?...)

180

Regressão Linear

Introdução

Como vimos anteriormente, um dos primeiros passos no estudo da relação entre duas variáveis métricas é elaborar uma nuvem de pontos, isto é, um gráfico que traduza a variação da variável dependente (Y) em função da variação da variável independente (X). Consideremos a nuvem de pontos representada na Figura abaixo que sugere existir uma tendência linear positiva. Se os coeficientes de correlação quantificam a relação, a **Análise de Regressão** é utilizada para modelar essa relação, isto é, construir uma equação que traduza a variação de y dado x.

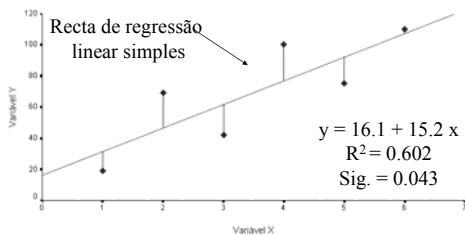


181

Regressão Linear

Introdução

Ao modelo atrás referido chama-se recta de regressão, representada na Figura abaixo. O critério de ajuste da recta de regressão consiste na minimização dos quadrados dos desvios da recta relativamente aos pontos observados. Por outras palavras, pretende-se obter a recta que torne mínimos os quadrados das distâncias verticais dos pontos observados à recta ajustada - erros (ver Figura abaixo). Na Regressão é testada, entre outros parâmetros, a significância do modelo estimado, e a qualidade do ajuste, dada pelo coeficiente de determinação - R^2 .



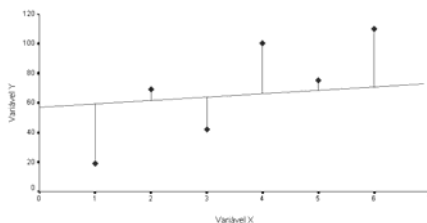
182

Regressão Linear

Introdução

Na Figura abaixo é representado um exemplo de um modelo cuja qualidade do ajuste é menor relativamente ao anterior, ou seja, os desvios verticais são maiores no modelo da Figura anterior. O modelo de Regressão linear simples (uma variável independente) possui dois parâmetros que definem a recta de regressão; a ordenada na origem (β_0 ou a) e o declive da recta (β_1 , b ou m). Utilizando a nomenclatura β_0 e β_1 , a equação da recta de regressão de y dado x é dada por:

$$y = \beta_0 + \beta_1 \cdot x$$



NOTA:
 A Regressão é uma técnica paramétrica que exige normalidade dos erros, entre outros pressupostos (Snedecor & Cochran, 1989).

183



Regressão Linear

Introdução

No cálculo dos parâmetros da regressão são utilizadas as seguintes fórmulas:

$$\beta_1 = \frac{S_{xy}}{S_{xx}} = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sum (x_i - \bar{x})^2} \quad \beta_0 = \bar{y} - \beta_1 \bar{x}$$

Para testar a significância da regressão procede-se a uma análise de variância (*):

Fonte de variação	Soma de quadrados (SQ)	gl's	Quadrados médios	F
Regressão	$\sum_{i=1}^n (\hat{y}_i - \bar{y})^2$	1	$\frac{\text{SQ regressão}}{\text{gl's regressão}}$	SQ regressão gl's regressão SQ erro gl's erro
Erro	$\sum_{i=1}^n (y_i - \hat{y}_i)^2$	n - 2	$\frac{\text{SQ erro}}{\text{gl's erro}}$	
Total	$\sum_{i=1}^n (y_i - \bar{y})^2$	n - 1		

A estatística calculada F distribui-se, sob H_0 , como um $F_{(\alpha, 1, n-2)}$.

184



Regressão Linear

Introdução

O coeficiente de determinação (R^2) é dado pela expressão abaixo. Note-se que, no caso da regressão linear simples, R^2 é o quadrado do coeficiente de correlação linear de Pearson entre as duas variáveis.

$$R^2 = \frac{\text{SQ regressão}}{\text{SQ total}} = \frac{\sum (\hat{y}_i - \bar{y})^2}{\sum (y_i - \bar{y})^2}$$

R^2 varia entre 0 e 1 (ajuste perfeito, isto é, a recta é ajustada sobre todos os pontos observados).

Procedimentos em SPSS

```

Analyze
  Regression
    Linear
      Dependent variável dependente
      Independent(s) variável independente
      Statistics estimates model fit Continue
      Options include constant in equation Continue
  OK

```

185



Regressão Linear

Exemplo em SPSS

Dados

Consideremos o exemplo já trabalhado: amostra de 8 crianças e adolescentes dos quais se registou a idade e a estatura (cm). Como objectivo, pretendemos criar um modelo que explique a variação da estatura (variável dependente - Y) em função da idade (variável independente - X).

	crianças	idade	estatura
1	1	14	155.00
2	2	6	120.00
3	3	11	145.00
4	4	7	123.00
5	5	13	161.00
6	6	12	150.00
7	7	2	91.00
8	8	4	103.00

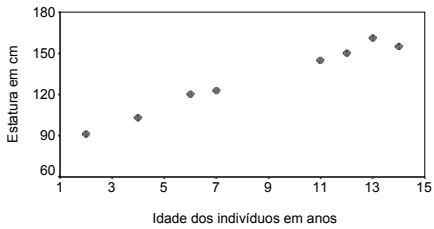
186

Regressão Linear

Exemplo em SPSS

Nuvem de pontos

O primeiro passo é sempre analisar a nuvem de pontos. Se a mesma não sugerir uma relação linear, o modelo de regressão linear não é, obviamente, adequado.

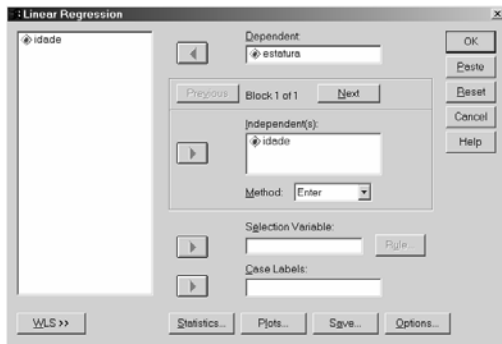


187

Regressão Linear

Exemplo em SPSS

Menu: Analyze \ Regression \ Linear

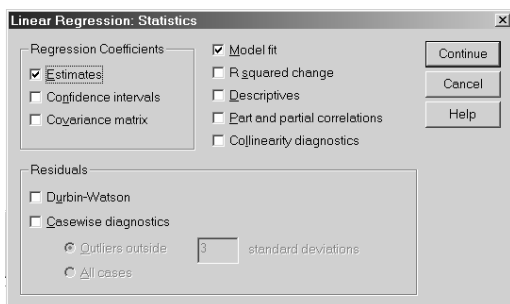


188

Regressão Linear

Exemplo em SPSS

Menu: Analyze \ Regression \ Linear \ Statistics



189

Regressão Linear

Análise Avançada de Dados

Exemplo em SPSS

Menu: Analyze \ Regression \ Linear \ Options

Linear Regression: Options

Stepping Method Criteria:

☒ Use probability of F
Entry: .05 Removal: .10

☐ Use F value
Entry: 3.84 Removal: 2.71

☒ Include constant in equation

Missing Values:

☒ Exclude cases listwise
☐ Exclude cases pairwise
☐ Replace with mean

Continue Cancel Help

190

Regressão Linear

Análise Avançada de Dados

Exemplo em SPSS

Resultados:

Resumo do modelo: Coeficiente de determinação

Model Summary

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate
1	.990 ^a	.979	.976	3.976

a. Predictors: (Constant), IDADE

O valor do coeficiente de determinação (*R-Square* = 0.979) é muito elevado - perto do valor máximo possível: 1. Deste modo existe uma boa qualidade de ajuste do modelo aos dados observados.

191

Regressão Linear

Análise Avançada de Dados

Exemplo em SPSS

Resultados: Quadro de ANOVA

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	4507.139	1	4507.139	285.077	.000 ^a
	Residual	94.861	6	15.810		
	Total	4602.000	7			

a. Predictors: (Constant), IDADE

b. Dependent Variable: ESTATURA

A regressão é significativa para um nível de significância de, p. ex., 0.05, pois *Sig.* < 0.000. Por outras palavras (relembrando...), a probabilidade de se cometer um erro tipo I (α), ou seja, rejeitar H_0 quando é verdadeira, é muito pequena.

192

Regressão Linear

Exemplo em SPSS

Resultados: Coeficientes

Já sabendo que a regressão é significativa, quais são os parâmetros do modelo que permitem definir uma equação que traduza a variação de Y dado X?

Coefficients^a

Model	Unstandardized Coefficients		Standardized Coefficients	t	Sig.
	B	Std. Error	Beta		
1	(Constant)	82.040	3.223	25.458	.000
	IDADE	5.676	.336	16.884	.000

a. Dependent Variable: ESTATURA

$B_0 = 82.04$; $B_1 = 5.68$. Então a equação toma a forma:

$$\text{Estatura} = 82.0 + 5.7 * \text{Idade}$$

193

Regressão Linear

Exemplo em SPSS

Resultados: Comentários gerais.

É importante relembrar que a utilização do modelo apenas é correcta no intervalo de valores de X com o qual foi elaborado: neste caso, 2 a 14 anos.

Com base nesta equação estimemos a estatura de uma criança com 8 anos de idade:

$$\text{Estatura com 8 anos} = 82.0 + 5.7 * 8$$

$$\text{Estatura com 8 anos} = 127.6 \text{ cm}$$

Em SPSS a linguagem **syntax** permite efectuar estes e outros cálculos. Façamos de seguida uma breve introdução a esta linguagem.

194

A linguagem syntax

Todas as acções efectuadas por menus em SPSS podem ser opcionalmente traduzidas em linguagem *syntax* no ficheiro *output*. Por exemplo, a syntax dum histograma de uma variável "var1" é:

```
GRAPH
  /HISTOGRAM=var1 .
```

Esta e outras "ordens" podem ser dadas / reproduzidas numa janela de syntax. Também a grande maioria das operações matemáticas e de cálculo podem ser efectuadas em syntax. Por exemplo, para se efectuar o cálculo anterior que nos dá a resposta do modelo a uma idade de 8 anos:

195

A linguagem syntax

Janela de syntax:

exemplo de aplicação do modelo linear estimado anteriormente



```

Syntax1 - SPSS Syntax Editor
File Edit View Analyze Graphs Utilities Run Window Help
[Icons]
matrix.
compute y=82.0+5.7*8.
print y /title="Altura com 8 anos".
end matrix.
SPSS Processor is ready
  
```

196

A linguagem syntax

Janela de syntax:

O que são os comandos atrás?

```
matrix.
PROCEDIMENTOS
end matrix.
```

As expressões `matrix` e `end matrix` são, como regra, colocadas ao início e ao fim do procedimento geral a efectuar.

```
compute y=FÓRMULA .
```

"compute y=" significa "calcula o valor de y igual a..."

```
print y /title="TÍTULO".
```

"print y" é o comando usado para visualizar / imprimir o valor calculado no ficheiro de *output* neste caso com a opção de título `"/title" "TÍTULO"`

Uma regra fundamental é colocar um ponto "." no final de cada linha ("ordem") definida.

197

A linguagem syntax

Janela de syntax:

Depois de escrito o pequeno programa que pede o cálculo de y dado x=8, faz-se *correr* o mesmo pelo menu **RUN \ ALL**. O resultado do nosso exemplo é:

```
matrix.
compute y=82.0+5.7*8.
print y /title="Altura com 8 anos".
end matrix.
```

Run MATRIX procedure:

```
Altura com 8 anos
127.6000000
```

```
----- END MATRIX -----
```

Nota final: A potencialidade e utilidade da linguagem syntax estende-se muito para além deste pequeno exemplo: o SPSS possui manuais de ajuda através do menu **HELP \ SYNTAX GUIDE \ BASE**.

198

Introdução à Amostragem

Enquadramento...

Como nota prévia é de referir que a amostragem consiste numa grande *sub-secção* da estatística, bastante bem documentada na literatura. Este capítulo apresenta apenas uma (muito) breve introdução a toda a problemática da amostragem que, por si só, mereceria um curso completo de um ou mais semestres.

Para futuras abordagens, ficam as referências:

Cochran, W., 1977. *Sampling Techniques*, 3rd. ed., Wiley. 428 pp.
(uma referência importante da amostragem)

Malhotra, K. and Birks, D., 2000. *Marketing Research, An Applied Approach*. 3rd. ed., Financial Times – Prentice-Hall. 731 pp.
(uma referência mais *ligeira*, aplicada ao marketing)

Introdução à Amostragem

Os passos indispensáveis na amostragem

Os conceitos “População”, “amostra” e “inferência” já foram previamente abordados. As etapas fundamentais num processo de amostragem consistem na definição dos seguintes pontos:

1. Objectivos do estudo;
2. População a amostrar;
3. Dados a recolher;
4. Grau de precisão;
5. Métodos de medida;
6. Unidades amostrais;
7. Selecção da amostra.

Introdução à Amostragem

O conceito de aleatoriedade.

A amostragem probabilística obedece a determinadas propriedades matemáticas* que permitem poder-se desenvolver *a posteriori* toda uma teoria de amostragem.

O conceito de aleatório tem a ver com conhecer-se *a priori* a probabilidade de ser seleccionada uma certa unidade amostral ou amostra. No caso mais elementar, quando esta probabilidade é igual para todas as unidades amostrais / amostras, estamos em presença da **amostragem aleatória simples** (voltaremos a este assunto posteriormente).

Introdução à Amostragem

? O que é uma boa amostra ?



- Representa o melhor possível as características da População.
- Formalmente, maximiza a PRECISÃO* e minimiza o ENVIESAMENTO*.

202

Introdução à Amostragem

ENVIESAMENTO (*bias*)



1. Porque é que existe enviesamento nos exs. abaixo?
2. Sugestões para melhorar o problema...

- Referendo à população portuguesa sobre autoestrada Braga – Vigo ou conclusão da Via do Infante (Algarve); população inquirida: Distrito do Porto.
- Um governo pretende saber o número de desempregados do país. Para isso, extrai de uma listagem geral uma amostra de 860 indivíduos com idades superiores a 12 anos.
- Percentagem da população portuguesa que é católica e praticante: amostra de 240 indivíduos da Beira Interior.

203

Introdução à Amostragem

PRECISÃO / RIGOR*



1. Porque é que existe imprecisão nos exs. abaixo?
2. Sugestões para melhorar o problema...

- A análise de uma amostra dos estudantes do ensino superior algarvio conclui que o tempo médio de deslocação à instituição de ensino é de 55 minutos, e que os mesmos demoram entre 2 minutos e 1h 52' na deslocação.

204

Introdução à Amostragem

Métodos de selecção da amostra

Os diferentes métodos de selecção da(s) amostra(s) estão associados ao modo como se trata as características da população alvo. Por exemplo, se a população é finita ou não^{*1}, ou se existe reposição^{*2} ou não (particularmente importante no caso de populações finitas).

Tendo em conta este tipo de condicionantes, referem-se abaixo alguns tipos de amostragem.

1. Aleatória simples (a.a.s.)
2. Aleatória estratificada
3. Por conglomerados
4. Sistemática

Exemplos...

205

Introdução à Amostragem

A.a.s. simples

Consiste em seleccionar n indivíduos (amostra) de entre N (população) de tal modo que todas as combinações possíveis de n possuam a mesma probabilidade de serem seleccionadas. Na prática o processo aleatório pode ser simulado de diferentes maneiras, a mais simples das quais através de uma tabela de números aleatórios*. Existem também algoritmos de geração de números aleatórios implementados em muitos programas de computador**.

Ex: a extração do totoloto consiste num processo de a.a.s. sem reposição onde uma amostra $n = 6$ é seleccionada de uma população finita $N = 49$ e no qual o processo aleatório é simulado mecanicamente.

Façamos as contas à probabilidade de ficarmos milionários: totoloto vs. lotaria (revisitado)...

206

Introdução à Amostragem

As características a observar

Num processo de amostragem decide-se qual ou quais as características a medir / observar nos indivíduos seleccionados: o peso, sexo, preferência por um dado produto, o número da bola no caso do totoloto, etc. No entanto, de modo geral, o interesse centra-se nas seguintes quatro características da população:

1. Média (p.ex., rendimento médio anual das pme's do Algarve)
2. Total (p.ex., número total de peixes num manancial)
3. Razão (p.ex., entre rendimentos líquidos e brutos das famílias de uma região)
4. Proporção (p.ex., proporção de fumadores masculinos de uma Universidade)

207

Introdução à Amostragem

Amostragem aleatória simples: os estimadores

No caso da a.a.s., os estimadores amostrais utilizados na inferência destas quatro características da população são:

$$\text{Média populacional } \bar{Y} \quad \hat{\bar{Y}} = \bar{y} = \text{média amostral}$$

$$\text{Total populacional } Y \quad \hat{Y} = N \bar{y} = N \frac{\sum_{i=1}^n y_i}{n}$$

$$\text{Razão populacional } R \quad \hat{R} = \frac{\bar{y}}{\bar{x}} = \frac{\sum_{i=1}^n y_i}{\sum_{i=1}^n x_i}$$

$$\text{Proporção populacional } P \quad \hat{P} = \frac{a}{n}$$

208

Introdução à Amostragem

O tamanho da amostra: definições

Os métodos de determinação do tamanho da amostra variam de acordo com as características do processo de amostragem. Analisemos dois casos simples: adicionemos as seguintes definições às atrás referidas:

- t desvio da curva Normal (abscissa) correspondente à confiança probabilística: p. ex., para um nível de confiança de 95%, $t = 1.96$; para 99%, $t = 2.58$.
- p estimativa da proporção de indivíduos numa classe C (por ex., fumadores).
- q estimativa da proporção de indivíduos na classe complementar C' (p.ex., não fumadores): $q = 1 - p$.

209

Introdução à Amostragem

O tamanho da amostra: definições - continuação

- d erro estabelecido *a priori*.
- s estimativa do desvio padrão da população.
- n_0 primeira aproximação ao tamanho amostral (fórmulas abaixo).
- n tamanho amostral após análise da fracção amostral: se a mesma fôr superior a 10%, corrigir n_0 através das fórmulas abaixo (correção para população finita - *cpf*); se $f < 10\%$, considera-se $n = n_0$.

210

Introdução à Amostragem

O tamanho da amostra:

Fórmulas

Proporções

$$n_0 = \frac{t^2 p q}{d^2}$$

Dados contínuos

$$n_0 = \frac{t^2 s^2}{d^2}$$

Se $f(n_0 / N) > 0.1$ (10%), corrigir com:

$$n = \frac{n_0}{1 + \frac{(n_0 - 1)}{N}}$$

Caso contrário tomar $n = n_0$

211

Introdução à Amostragem

Estimação do tamanho da amostra - 2 exemplos:

1. Pretende-se estimar n associado à estimativa $p = 0.35$, sendo P = proporção de indivíduos fumadores do INUAF, adoptando um erro de 0.05, um nível de confiança de 95% ($t = 1.96$), e $N = 2000$.

$$n_0 = \frac{1.96^2 0.35(1 - 0.35)}{0.05^2} \cong 350, \text{ sendo } f = \frac{350}{2000} \cong 0.17$$

Uma vez que a fracção amostral f supera 0.1 procede-se à correcção para pop. finita:

$$n = \frac{350}{1 + \frac{(350 - 1)}{2000}} = 298$$

212

Introdução à Amostragem

Estimação do tamanho da amostra - 2 exemplos:

2. Qual o tamanho da amostra quando se pretende estimar o rendimento médio mensal por agregado familiar no concelho de Loulé? Considere 700 € como estimativa do desvio padrão da população, 100 € como o erro admissível, 99% como nível de confiança ($t = 2.58$) e 20 000 habitantes como a estimativa da população do concelho.

$$n_0 = \frac{2.58^2 700^2}{100^2} = 326, \text{ sendo } f = \frac{326}{20000} \cong 0.02$$

Uma vez que a fracção amostral f é menor que 0.1, $n = n_0 = 326$

213



Introdução à Amostragem

Tamanho da amostra:

Notas finais:

1. Em muitas situações pretende-se a estimação de mais que um parâmetro. Por exemplo, um questionário com múltiplas perguntas. Neste caso há que estimar o tamanho amostral separadamente para cada um dos parâmetros mais importantes e **seleccionar o maior** destes valores.
2. Existem muitas outras abordagens na estimação de n , não só dependentes do tipo de amostragem efectuado mas também resultantes da incorporação de outros parâmetros na estimação como por exemplo o custo da amostragem (ver, para este efeito, Cochran, 1977).
