

Estudo das propriedades psicométricas das escalas

A validação da recolha de informações é o processo pelo qual o investigador ou o avaliador se assegura que aquilo que quer recolher como informações, as informações que recolhe realmente e o modo como as recolhe servem adequadamente o objectivo da investigação.

Previamente a qualquer tratamento estatístico, procedeu-se à observação geral dos dados constantes na base, de forma a evitar erros não somente no carregamento da base (dados em falta) como a eliminação de todos os dados para os quais não existia correspondência entre a resposta da chefia e resposta do colaborador, findo o qual se deu início à análise dos dados propriamente dita.

Para avaliar as qualidades psicométricas de qualquer instrumento de medida, é necessário efectuar estudos de **fidelidade** e **validade**, que nos indicam o grau de confiança nas escalas e de generalização que os dados poderão alcançar.

Assim os estudos de fidelidade fornecem indicações sobre o grau de confiança ou exactidão que podemos ter na informação obtida, ou seja a sua consistência, enquanto o estudo da validade se refere à avaliação do grau em que uma determinada medida mede, de facto, o que se pretende medir.

Por outro lado, o estudo das propriedades da distribuição normal é extremamente útil e importante quando o investigador pretende fazer inferências sobre a população e só dispõe de dados referentes a uma amostra.

- **FIDELIDADE**

O objectivo de qualquer instrumento é avaliar as características dos sujeitos, não podendo por tal ser ambíguo ou originar diferentes interpretações, mas sim ser interpretado de forma similar por todos.

O *alpha* de Cronbach trata-se de uma das medidas vulgarmente utilizadas para verificação da fidelidade interna, a qual traduz essencialmente a média de todos os coeficientes de bi-partição possíveis.

O *alpha* de Cronbach é tradicionalmente utilizado em escalas tipo Likert sendo apontado como o indicador mais importante de fiabilidade de um instrumento.

A determinação deste coeficiente tanto para o total da escala como para as dimensões, permite estimar a homogeneidade dos itens, isto é, até que ponto cada enunciado da escala mede o mesmo conceito de forma equivalente.

De modo a facilitar a interpretação, apresenta-se na Tabela 1 a relação entre o *alpha* de Cronbach e a qualidade da fidelidade da escala.

Tabela 1 - Fidelidade de uma escala

Alpha de Cronbach	Fidelidade
> 0,9	Excelente
0,8 – 0,9	Boa
0,7 – 0,8	Razoável
0,6 – 0,7	Fraca
<0,6	Inaceitável

Fonte: Hill e Hill (2002)

 **Instruções e resultados em SPSS:**

SPSS:  Analyse  Scale  Reability Analysis ...

Exemplo:

Consideremos a seguinte codificação das respostas do questionário:

NADA –1 ; POUCO - 2 ; MUITO – 3 ; BASTANTE – 4.

1,00	2,00	1,00	2,00	2,00	1,00	1,00	1,00	1,00	1,00
3,00	4,00	3,00	3,00	4,00	3,00	3,00	4,00	3,00	3,00
2,00	2,00	2,00	2,00	2,00	2,00	2,00	3,00	2,00	3,00
1,00	2,00	1,00	2,00	2,00	1,00	1,00	1,00	1,00	1,00
4,00	4,00	4,00	4,00	4,00	3,00	3,00	4,00	4,00	4,00
2,00	2,00	2,00	1,00	1,00	2,00	2,00	1,00	2,00	1,00
3,00	4,00	3,00	3,00	4,00	2,00	4,00	3,00	2,00	3,00
2,00	3,00	2,00	2,00	3,00	3,00	3,00	3,00	2,00	3,00
1,00	1,00	1,00	1,00	1,00	1,00	1,00	1,00	1,00	1,00
4,00	4,00	4,00	4,00	4,00	3,00	3,00	3,00	4,00	4,00
1,00	2,00	2,00	2,00	1,00	1,00	1,00	2,00	2,00	2,00
3,00	3,00	3,00	3,00	3,00	4,00	4,00	3,00	4,00	3,00
1,00	1,00	2,00	3,00	2,00	3,00	2,00	1,00	2,00	3,00
2,00	2,00	2,00	2,00	2,00	2,00	1,00	1,00	2,00	1,00
2,00	3,00	3,00	4,00	2,00	1,00	3,00	4,00	3,00	1,00

Reliability Statistics		Item-Total Statistics				
		Scale Mean if Item Deleted	Scale Variance if Item Deleted	Corrected Item-Total Correlation	Cronbach's Alpha if Item Deleted	
Cronbach's Alpha	N of Items	P1	21,27	69,495	,909	,955
		P2	20,80	70,314	,862	,957
		P3	21,07	70,495	,930	,955
		P4	20,87	72,552	,780	,960
		P5	20,93	69,210	,866	,957
		P6	21,27	74,067	,683	,964
		P7	21,13	69,981	,842	,958
		P8	21,07	68,638	,808	,960
		P9	21,07	70,352	,868	,957
		P10	21,13	69,981	,789	,960
		,962	10			

Pode observar-se que a consistência interna deste questionário é excelente, o valor de alpha é 0,962.

- **VALIDADE**

A validade refere-se à avaliação do grau em que uma determinada medida mede, de facto, o que se pretende medir.

O estudo da validade, utilizando a análise factorial dos itens e dos seus resultados, é o método que tem revelado maior uso e reconhecimento entre os diversos autores.

A análise factorial é uma técnica de análise exploratória de dados que tem por objectivo descobrir e analisar a estrutura de um conjunto de variáveis interrelacionadas de modo a construir uma escala de medida para factores que de alguma forma controlam as variáveis originais (Maroco, 2007). Com o objectivo da redução da dimensão das variáveis, sem perda de informação, a análise factorial foi elaborada com vista a identificar um conjunto menor de variáveis hipotéticas (componentes). A rotação *Varimax* faz com que, para cada componente principal, existam apenas alguns pesos significativos e todos os outros sejam próximos de zero.

A análise factorial apenas tem utilidade na estimação de factores comuns, quando as correlações entre as variáveis originais são elevadas o suficiente. O teste de Esfericidade de Barlett pode ser utilizado, contudo, é muito sensível à dimensão da amostra. O método de utilização mais popular é a “medida da adequação da amostragem de Kaiser-Meyer-Olkin”, a qual mede a homogeneidade das variáveis, cujas recomendações constam da Tabela 2.

Tabela 2 - Valor do KMO e recomendação relativamente à Análise Factorial

KMO	Análise Factorial
0,9 – 1,0	Excelente
0,8 – 0,9	Boa
0,7 – 0,8	Média
0,6 – 0,7	Medíocre
0,5 - 0,6	Mau, mas ainda aceitável
<0,5	Inaceitável

Fonte: Maroco, 2007

○ APRESENTAÇÃO DO MÉTODO.

Suponha agora que pretende avaliar o grau de satisfação pessoal de cada indivíduo, para tal elaborou um questionário em que num determinado item se pergunta qual o grau de satisfação com a sua profissão (item 1) e noutro item se pergunta qual o grau de satisfação obtido com o seu trabalho diário (item 2). Quase de certeza os dois itens estarão muito correlacionados. Se a correlação entre estes dois itens for elevada poderia afirmar-se que são redundantes, isto é, no questionário bastaria utilizar apenas um deles. Se conseguirmos determinar uma nova variável que traduza a relação entre estas variáveis (item 1 e item 2), num futuro questionário as respostas a esta nova variável representariam de modo idêntico as respostas aos dois itens iniciais. Estatisticamente combinámos duas variáveis num factor que traduz a “essência” dessas variáveis (o novo factor é uma combinação linear das duas variáveis).

Este exemplo ilustra a **ideia básica** da *Análise de Componentes Principais (ACP)*:

- *reduzir o número de variáveis iniciais* através da procura de factores que traduzam a “essência” dessas variáveis.
- *detectar uma estrutura na relação entre as variáveis iniciais*, ou seja, classificá-las em grupos semelhantes.

Suponha que dispõe de $X_1, \dots, X_k$ variáveis iniciais de escala métrica, a ACP é uma das técnicas de Análise de Dados que nos permite determinar novas variáveis (factores) que traduzem a mesma variabilidade (informação) que o conjunto inicial de variáveis.

○ DETERMINAÇÃO DAS COMPONENTES, EIXOS E FACTORES PRINCIPAIS

Considere que dispõe de k variáveis medidas sobre n indivíduos, o que constitui a matriz inicial de dados:

Variáveis indivíduos	X_1	...	X_j	...	X_k
1	X_{11}	...	X_{j1}	...	X_{k1}
2	X_{12}	...	X_{j2}	...	X_{k2}
...	...	...	...	...	...
i	X_{1i}	...	X_{ji}	...	X_{ki}
...	...	...	...	...	...
n	X_{1n}	...	X_{jn}	...	X_{kn}

Sendo X_{ji} a resposta ao item j (variável j) do indivíduo i, $j=1,\dots,k$; $i=1,\dots,n$.

Geométricamente extrair as componentes principais equivale a proceder a uma rotação do espaço gerado pelas variáveis iniciais $X_1,\dots,X_k$ (que frequentemente não é ortogonal porque as variáveis possuem correlações não nulas entre si) para um novo espaço constituído por eixos ortogonais (os factores). Este tipo de rotação chama-se “**variance maximizing**” (maximização da variância) porque o critério de rotação do sistema de eixos iniciais consiste na maximização da variância da “nova” variável (factor), enquanto minimiza a variância entre esta “nova” variável (1º factor) e a seguinte (2º factor), etc..

As variáveis iniciais, cada uma com variância S_j^2 , $j=1,\dots,k$, apresentam uma variabilidade total $\sum_{j=1}^k S_j^2$, mas como usualmente as variáveis iniciais são normalizadas pelo package estatístico, e a variância de cada variável normalizada é a unidade, a variabilidade total é k, isto é, o número de variáveis iniciais.

Em ACP após a extração do primeiro eixo/factor (o que possui maior variância) procura-se um segundo eixo/factor que maximize a restante variância e assim sucessivamente. Por este processo extraem-se os k factores, e uma vez que cada factor está definido de modo a maximizar a variabilidade total que não foi “capturada” pelo factor precedente, os **factores** assim obtidos são **independentes**, isto é, não relacionados ou seja **ortogonais**. Então poderemos afirmar que de um sistema de eixos gerado pelas variáveis iniciais obtemos um novo sistema de eixos ortogonais (a informação fornecida por um eixo é diferente da informação fornecida por outro eixo). A este novo sistema de eixos ortogonais chamamos **factores principais**.

Teoricamente extrair as k componentes principais que definem a direcção dos k eixos ortogonais (factores principais) equivale a determinar os k valores próprios (eigenvalues) $\lambda_1, \dots, \lambda_k$ e os respectivos vectores próprios $w_1, \dots, w_k$, associados à matriz de correlações R das variáveis iniciais:

$$R = \begin{bmatrix} 1 & r(X_1, X_2) & \dots & r(X_1, X_j) & \dots & r(X_1, X_k) \\ & 1 & \dots & r(X_2, X_j) & \dots & r(X_2, X_k) \\ & & \dots & \dots & \dots & \dots \\ & & & 1 & \dots & r(X_j, X_k) \\ & & & & \dots & \dots \\ & & & & & 1 \end{bmatrix}$$

Os vectores próprios desta matriz R : $w_1 = \begin{bmatrix} w_{11} \\ w_{21} \\ \dots \\ w_{j1} \\ \dots \\ w_{k1} \end{bmatrix}, \dots, w_i = \begin{bmatrix} w_{1i} \\ w_{2i} \\ \dots \\ w_{ji} \\ \dots \\ w_{ki} \end{bmatrix}, \dots, w_k = \begin{bmatrix} w_{1k} \\ w_{2k} \\ \dots \\ w_{jk} \\ \dots \\ w_{kk} \end{bmatrix}$, são os

pesos de cada variável nas várias componentes.

Os valores próprios da matriz R são as variâncias de cada componente, isto é, $\text{var}(C_1) = \lambda_1 > \text{var}(C_2) = \lambda_2 > \dots > \text{var}(C_k) = \lambda_k$, e como referimos anteriormente a variância da 1ª componente é maior do que da 2ª componente, ..., etc. É de notar que ao extrair os k factores principais, as respectivas direcções fornecidas pelas *componentes principais mantêm a informação fornecida pelas variáveis iniciais*, já que se as variáveis são normalizadas se tem uma variabilidade total igual a k, e prova-se que também a variabilidade das k componentes é $\sum_{j=1}^k \lambda_j = k$. O modo mais comum de apresentar os valores próprios é em termos de percentagem de variância total explicada:

Valores próprios (variância das componentes)	% explicada pela componente	% acumulada
λ_1	λ_1/k	λ_1/k
λ_2	λ_2/k	$(\lambda_1 + \lambda_2)/k$
...	...	...
λ_k	λ_k/k	$(\lambda_1 + \lambda_2 + \dots + \lambda_k)/k$

Em ACP quantos factores devemos extrair?

É de notar que associados a k variáveis iniciais há também k factores. Mas se formos extrair k factores não estamos a reduzir a dimensionalidade do espaço de trabalho, como é do objectivo de uma ACP, apenas ganhamos no tipo de informação (mais interpretável, pelo menos teoricamente) fornecida pelas k componentes dado que como referimos anteriormente elas são ortogonais entre si. Temos também interesse em escolher das k componentes apenas um número restrito delas que não explicando o total da informação fornecida pelos dados (variabilidade das k variáveis iniciais), nos explique pelo menos a maioria dessa informação (o que podemos observar através da % acumulada de variância das componentes extraídas).

A questão agora é quantos factores principais devemos extrair? Note que o processo de extração dos vários factores é realizado de modo a que cada factor tenha sucessivamente menos variabilidade. A regra de paragem depende da percentagem de variabilidade que se pretende deixar por explicar. Esta decisão é portanto arbitrária e depende não só do utilizador, como também do tipo de estudo. Há no entanto alguns critérios de decisão dos quais o mais utilizado é o ***critério de Kaiser***.

Segundo o critério de Kaiser devemos *reter apenas os factores cuja variância seja superior a um*. Essencialmente a ideia básica desta escolha consiste em reter os factores cuja variância seja superior à de cada variável inicial (normalizada), caso contrário a informação fornecida por esse factor é inferior à de cada variável inicial.

○ OBTENÇÃO DOS RESULTADOS E SUA INTERPRETAÇÃO

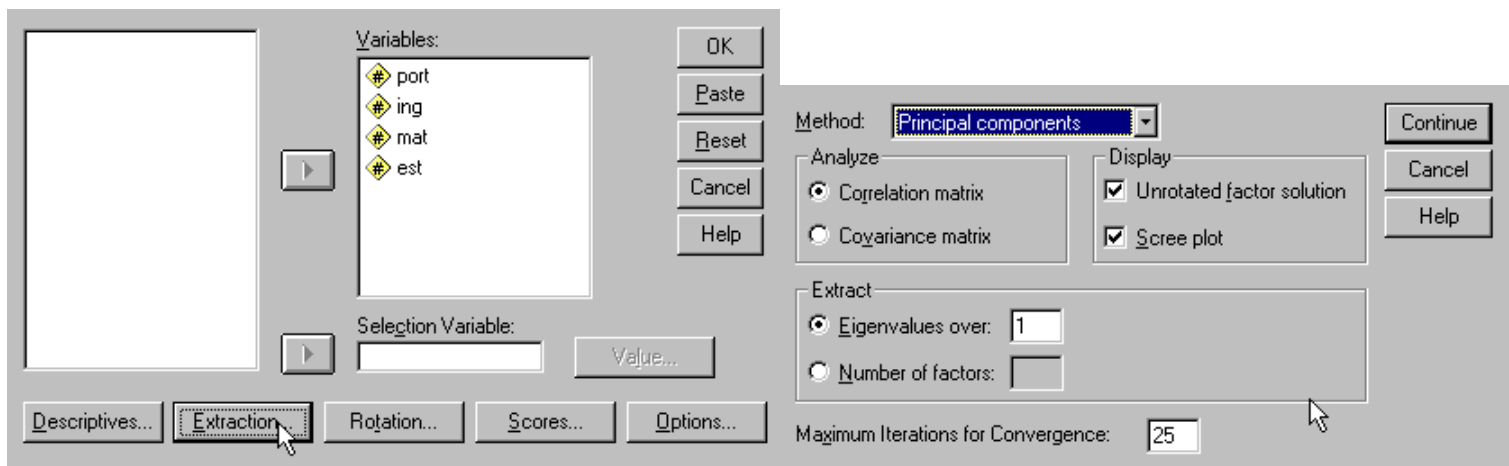
Considere o seguinte exemplo: foram registadas as pontuações de seis alunos nas disciplinas de Português, Inglês, Matemática e Estatística, e obteve-se a seguinte matriz de dados:

Alunos	Português	Inglês	Matemática	Estatística
A	15	14	10	12
B	14	13	11	13
C	16	15	10	11
D	18	16	11	12
E	12	12	10	10
F	15	14	13	14

 Instruções e resultados em SPSS:

SPSS:  Analyse  Data Reduction  Factor

Na opção Descriptives seleccionar KMO and Bartlett's test of sphericity



KMO and Bartlett's Test

Kaiser-Meyer-Olkin Measure of Sampling Adequacy.		,410
Bartlett's Test of Sphericity	Approx. Chi-Square	16,529
	df	6
	Sig.	,011

Observando o valor de KMO=0,410 será inaceitável fazer uma ACP.

No teste de Bartlett, $H_0: \Sigma = I$, ora se se rejeitar a hipótese nula, tem sentido fazer ACP, já que se pode afirmar que existem variáveis relacionadas entre si, pelo que se podem constituir novas variáveis, os factores, formados por grupos de variáveis iniciais. Neste caso rejeita-se a hipótese nula do teste de Bartlett, já que $\text{sig}=0,011 < \alpha=0,05$.

Como podemos observar pela figura o SPSS utiliza o critério de Kaiser para seleccionar os factores principais (eigenvalues over 1.00).

Total Variance Explained

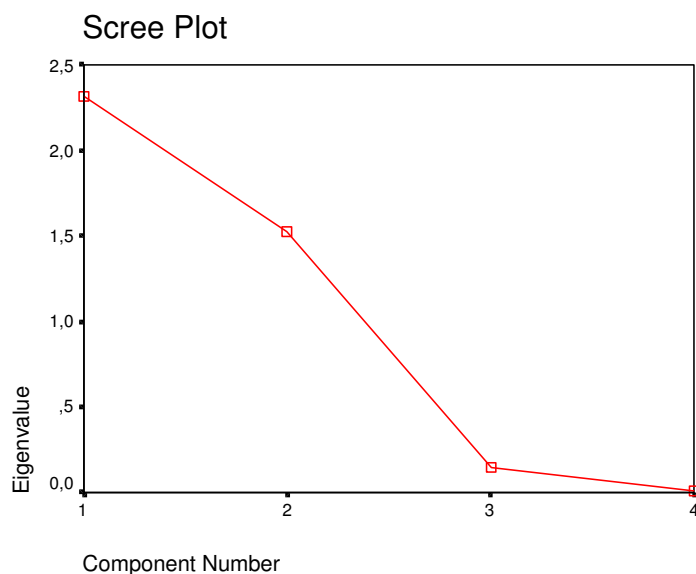
Component	Initial Eigenvalues			Extraction Sums of Squared Loadings		
	Total	% of Variance	Cumulative %	Total	% of Variance	Cumulative %
1	2,316	57,892	57,892	2,316	57,892	57,892
2	1,529	38,219	96,110	1,529	38,219	96,110
3	,150	3,752	99,862			
4	5,509E-03	,138	100,000			

Extraction Method: Principal Component Analysis.

Com a saída “Total Variance Explained” podemos observar a variância de cada componente principal extraída na análise e a percentagem de variância explicada por estas componentes:

isto é, com a 1ª componente explicamos 57,892% da informação fornecida pelos dados e com a 2ª componente 38,219% , ou seja, com apenas estas duas componentes explicamos 96,11% do total da informação.

Com o scree plot podemos observar a evolução dos valores próprios extraídos:



Component Matrix^a

	Component	
	1	2
PORT	,848	-,528
ING	,807	-,586
EST	,723	,634
MAT	,650	,711

Extraction Method: Principal Component Analysis.

a. 2 components extracted.

A próxima saída com interesse para a interpretação dos resultados é “component matrix”:

A saída component matrix fornece as correlações entre cada variável em estudo e o respectivo factor.

Se uma variável tiver uma grande correlação com o factor isto significa que essa variável é importante para definir o factor.

Neste exemplo atendendo às correlações de cada variável com o 1º factor podemos afirmar que este factor tem carácter genérico de factor comum entre as várias disciplinas, contudo, por serem maiores as correlações com as disciplinas de Português e de Inglês poderemos afirmar que o 1º factor é essencialmente o factor “LETRAS”. No 2º factor podemos observar que existe uma maior correlação das disciplinas de Matemática e Estatística com este factor, pelo que podemos afirmar ser um factor “CIÊNCIAS”.

Como podemos observar não existem regras fixas para interpretar cada uma das componentes principais (factores principais), apenas podemos realizar afirmações discutíveis sobre o conteúdo dessas componentes através dos valores das correlações de cada uma das variáveis com a componente. Isto é, podemos afirmar que uma determinada componente é constituída por um conjunto de variáveis iniciais, que apresentam grande semelhança entre si e grande correlação com o factor em causa.

Com a saída Communalities podemos extrair várias conclusões:

Communalities

	Initial	Extraction
PORT	1,000	,997
ING	1,000	,995
EST	1,000	,925
MAT	1,000	,927

Extraction Method: Principal Component Analysis.

1. Na segunda coluna, EXTRACTION, figura a percentagem de explicação de cada variável através da análise realizada, isto é, a percentagem de variância **acumulada** de cada variável explicada **até** ao número de factores extraídos. No exemplo, a variável pontuação em Português é explicada a 99,7% pelos dois primeiros factores.
2. Define-se comunalidade como a proporção de variância de uma determinada variável que é comum com as restantes variáveis em análise. Um estimador da comunalidade de uma variável é o valor de R^2 múltiplo.

OBSERVAÇÕES:

É necessário ter alguns cuidados na interpretação dos resultados desta Análise. Em primeiro lugar convém verificar se cada variável consegue uma percentagem de explicação razoável com a análise em curso (através da última coluna relativa aos factores extraídos, da saída communalities). Se uma variável não estiver “bem” explicada na análise, isto é, não tiver uma percentagem de explicação superior a 50% (no mínimo) no conjunto dos factores extraídos, mais correctamente esta variável não deveria ser incluída na análise e deveria figurar como ***variável suplementar***, ou então ter-se-ia de extrair mais um factor.

○ ROTAÇÃO DE EIXOS

Sempre que o 1º factor aparece como um factor genérico deve proceder-se a uma rotação do sistema de eixos de modo a que se distinga mais facilmente a importância de cada variável para a explicação dos eixos. Existem vários métodos de rotação de eixos, sendo que o mais utilizado é o método VARIMAX.

Assim, procedendo a uma rotação VARIMAX dos dados em estudo obtem-se:

Rotated Component Matrix^a

	Component	
	1	2
PORT	,989	,136
ING	,995	6,568E-02
EST	,151	,950
MAT	4,549E-02	,962

Extraction Method: Principal Component Analysis.

Rotation Method: Varimax with Kaiser Normalization.

a. Rotation converged in 3 iterations.

As disciplinas de Português e Inglês têm uma importância mais notória para a formação do 1º factor, enquanto que as disciplinas de Matemática e de Estatística aparecem com maior importância no 2º factor.

- **SENSIBILIDADE**

A sensibilidade representa a capacidade que um teste tem em discriminar os sujeitos, segundo o factor que está a ser avaliado, isto é, capacidade para fornecer respostas diferentes consoante os sujeitos da aplicação, sendo considerada sensível quando se assemelha à distribuição normal.

Um dos testes estatísticos mais usualmente utilizados para testar o ajustamento de distribuições amostrais a determinadas funções de distribuição teórica é o teste de Kolmogorov-Smirnov (K-S) (Maroco, 2007).

Segundo Maroco (2007) os métodos paramétricos são robustos à violação do pressuposto da normalidade desde que as distribuições não sejam extremamente enviesadas ou achatadas e que as dimensões das amostras não sejam extremamente pequenas.

Deste modo, os dados são analisados não somente em função do K-S, mas também das medidas de assimetria (*skewness*) e de achatamento (*kurtosis*).

- **TESTE DE AJUSTAMENTO DE KOLMOGOROV-SMIRNOV**

Um teste de ajustamento é um teste não-paramétrico, que serve para averiguar se uma amostra pode ser considerada como proveniente de uma certa distribuição. Tem particular interesse o teste de ajustamento à normal.

Pressupostos exigidos no teste de ajustamento de Kolmogorov-Smirnov:

- A amostra provém de uma distribuição contínua.
- Os parâmetros da distribuição em teste são pré-especificados e não deveriam ser estimados a partir da amostra.

No caso de se realizar um ajustamento à normal tem-se como hipótese nula:

$$H_0: X \sim N(\mu, \sigma)$$

Pretende-se testar se a amostra da variável conhecimento é proveniente de uma população com distribuição normal.

Quando se utiliza o teste de Kolmogorov-Smirnov estimando os parâmetros a partir da amostra viola-se um dos pressupostos, pelo que Lillefors efectuou uma

correção ao teste de Kolmogorov-Smirnov no caso de se proceder a um ajustamento à normal, efectuando uma estimação de parâmetros.

 Instruções e resultados em SPSS:

SPSS:  Analyse  Explore  Plots  Normality plots with tests

Tests of Normality

	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
conhecimento	,141	38	<u>056</u>	,922	38	,011

a. Lilliefors Significance Correction

Neste caso não se rejeita a hipótese de que a amostra seja proveniente de uma população normal ($\text{sig}=0,056 > \alpha=0,05$).