

6 Análise Estatística

Após a coleta dos dados, foi usada a análise estatística para alcançar três objetivos propostos anteriormente. O primeiro foi o de atestar se a escala F-PEC tem um bom enquadramento na indústria brasileira. A análise fatorial resultou nos três fatores referentes às três subescalas propostas por Astrachan et al. (2002), dando a segurança de que a escala é realmente empregável. Em seguida, a análise de clusters encontrou os tipos comuns na amostra e, por fim, a análise de correlação mostrou que pode haver uma correlação entre o grupo a que uma empresa pertence e algumas variáveis estudadas.

6.1. Dados Coletados

A pesquisa retornou 55 casos, dos quais 4 foram eliminados por não representarem casos factíveis de serem encontrados na população, restando 51 casos. A Tabela 43 no Apêndice B mostra os 51 casos, preservando o nome das empresas.

A partir dos dados coletados, os valores de cada variável foram calculados. A subescala Poder gerou 3 variáveis, nomeadas P1, P2 e P3. A variável P1 foi calculada de acordo com as respostas à questão 1a e 1b. O valor na escala Likert do percentual de propriedade direta membros familiares (F) foi dividido pelo somatório de todos os valores de propriedade, ou seja, de membros familiares (F), membros não familiares (N) e da holding (H). Desta forma, encontramos o valor de propriedade direta da família ($F_d = F / (F + N + H)$). Caso existam cotas de propriedade de uma holding, precisamos encontrar o valor de influência indireta da família pela holding (F_i). Esse valor é calculado a partir do percentual de cotas da família na companhia holding sobre a soma total de cotas das família e não familiares, ou seja, $F_i = (H_f / (H_f + H_n)) \times (H / (F + N + H))$. Somando F_d com F_i , calculou-se o valor da variável P1.

$$P1 = \frac{DF}{DF + DN + DH} + \frac{DH}{DF + DN + DH} \times \frac{HF}{HF + HN}$$

Onde:

DF = Controle direto de membros da família;

DN = Controle direto de membros fora da família, o não familiares;

DH = Controle direto da *Holding*;

HF = Controle da *Holding* por membros da família;

HN = Controle da *Holding* por membros fora da família.

As variáveis P2 e P3 receberam seus valores da mesma forma. Calculou-se o valor influência direta pelo valor da escala referente à faixa do percentual de cotas de membros da família sobre o somatório dos valores de cotas de familiares (F), não familiares (N) e indicados pela família (I). Então, $Fd = F / (F + N + I)$. A influência indireta foi calculada usando-se o coeficiente proposto por Klein et al. (2005) de 10% (0,1). Sendo assim, $Fi = (0,1) * I / (F + N + I)$. Somando-se os dois valores de influência (Fd e Fi), chegou-se aos valores das variáveis P2 e P3.

$$P2 \text{ e } P3 = \frac{F}{F + N + I} + 0,1 \times \frac{I}{F + N + I}$$

Onde:

F = Membros da família;

N = Membros fora da família, o não familiares;

I = Membros indicados pela família.

A subescala Experiência deu origem às quatro variáveis E1, E2, E3 e E4. As três primeiras foram calculadas da mesma maneira. Utilizando a curva de influência proposta por Astrachan et al. (2002), a primeira geração possui influência 0. A passagem da primeira geração para a segunda implica no maior grau de influência adicionada do que nas gerações seguintes. À medida que as gerações vão se sucedendo, o grau de influência adicionada vai caindo cada vez mais, de forma exponencial. Dessa maneira, segundo a curva, a primeira geração possui influência 0; a segunda geração possui influência 0,5; a terceira possui influência 0,5+0,52; a quarta possui influência de 0,5+0,52+0,53, e assim por diante.

A variável E4 recebe seu valor de acordo com a escala de intenção de perpetuidade a seguir:

Intenção de perpetuidade	Somente na geração atual	Até a geração de meus filhos	Até a geração de meus netos	Até a geração de meus bisnetos	De maneira perpétua
Valor da variável E4	1	2	3	4	5

Tabela 3: Escala de intenção de perpetuidade.

A subescala Cultura dá origem a 12 variáveis C1 até C12, que recebem seus valores diretamente das escalas Likert das questões correspondentes A até L, respectivamente.

A Tabela 39 do Apêndice B, relaciona as variáveis e seus respectivos valores.

6.2. Tratamento dos *Missing Values*

Em praticamente todas as questões contidas no questionário enviado aos participantes, o respondente tinha a opção de não respondê-las, caso não se aplicassem a seu negócio. Pensando nisso, quase todos os *missing values* foram tratados como resposta nula, ou seja, valor zero. As únicas exceções foram quanto às variáveis P2 e P3, que indicam o poder de influência da família no Conselho de Administração e Diretoria, respectivamente. Levando-se em consideração que os papéis dos dois conselhos são de extrema importância e presentes em qualquer empresa, mesmo que um conselho não exista de maneira formal, as decisões devem ser tomadas. Dessa forma, se um dos conselhos não existir, o outro assume seu papel e, por consequência, as variáveis possuem o mesmo valor. No caso de nenhum dos dois conselhos existirem, as variáveis P2 e P3 assumem o mesmo valor da variável P1, pois neste caso os proprietários são responsáveis por todas as decisões.

Sendo assim, a Tabela 40 do Apêndice B mostra valores com os *missing values* tratados.

6.3. Teste de Normalidade

Antes dos valores coletados serem transformados em uma escala normal, foi realizado um teste de normalidade com as variáveis para determinar sua aderência à curva normal. Devido ao tamanho reduzido de casos, escolheu-se o teste Shapiro Wilk. Todos os valores indicaram que a distribuição da população se assemelha a uma normal.

Tests of Normality						
	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
Zscore: Poder 1	,303	51	,000	,730	51	,000
Zscore: Poder 2	,203	51	,000	,884	51	,000
Zscore: Poder 3	,162	51	,002	,898	51	,000
Zscore: Experiência 1	,313	51	,000	,748	51	,000
Zscore: Experiência 2	,310	51	,000	,793	51	,000
Zscore: Experiência 3	,298	51	,000	,771	51	,000
Zscore: Experiência 4	,238	51	,000	,833	51	,000
Zscore: Cultura 1	,347	51	,000	,726	51	,000
Zscore: Cultura 2	,207	51	,000	,812	51	,000
Zscore: Cultura 3	,210	51	,000	,810	51	,000
Zscore: Cultura 4	,318	51	,000	,788	51	,000
Zscore: Cultura 5	,319	51	,000	,694	51	,000
Zscore: Cultura 6	,330	51	,000	,623	51	,000
Zscore: Cultura 7	,280	51	,000	,717	51	,000
Zscore: Cultura 8	,280	51	,000	,824	51	,000
Zscore: Cultura 9	,304	51	,000	,767	51	,000
Zscore: Cultura 10	,335	51	,000	,628	51	,000
Zscore: Cultura 11	,371	51	,000	,683	51	,000
Zscore: Cultura 12	,343	51	,000	,661	51	,000

a. Lilliefors Significance Correction

Tabela 4: Resultado do teste de normalidade executado no SPSS®.

6.4. Transformação dos Valores

Para que pudessem ser comparados em uma escala comum, os valores de todas as variáveis foram convertidos em sua transformada z. O resultado do procedimento realizado no software estatístico SPSS® está descrito a seguir, na Tabela 5. As novas variáveis contendo a transformada z das variáveis originais foram nomeadas ZP1, ZP2...ZC12, conforme na Tabela 41, no apêndice B.

<i>Descriptive Statistics</i>						
	<i>N</i>	<i>Minimum</i>	<i>Maximum</i>	<i>Mean</i>	<i>Std. Deviation</i>	<i>Variance</i>
Poder 1	51	0	1	,83	,243	,059
Poder 2	51	0	1	,69	,296	,087
Poder 3	51	0	1	,57	,353	,124
Experiência 1	51	0	1	,29	,288	,083
Experiência 2	51	0	1	,36	,300	,090
Experiência 3	51	0	1	,31	,302	,091
Experiência 4	51	0	5	3,08	1,885	3,554
Cultura 1	51	0	5	3,90	1,565	2,450
Cultura 2	51	0	5	3,37	1,766	3,118
Cultura 3	51	0	5	3,59	1,590	2,527
Cultura 4	51	0	5	3,65	1,508	2,273
Cultura 5	51	0	5	3,94	1,529	2,336
Cultura 6	51	0	5	4,12	1,451	2,106
Cultura 7	51	0	5	3,86	1,562	2,441
Cultura 8	51	0	5	3,45	1,616	2,613
Cultura 9	51	0	5	3,69	1,594	2,540
Cultura 10	51	0	5	4,10	1,418	2,010
Cultura 11	51	0	5	3,82	1,493	2,228
Cultura 12	51	0	5	4,00	1,400	1,960
Valid N (listwise)	51					

Tabela 5: Resultado da análise descritiva extraída do SPSS®.

6.5. Procura por *Outliers*

À procura de possíveis outliers, a tabela de transformada z das variáveis foi analisada. Valores de z-score com módulo superior a 2,5 devem ser vistos com possíveis outliers. Explorando os valores extremos de cada variável, nota-se que os casos 5, 7, 33, 19, 18, 11, 4 e 12 têm valores modulares de z-score superiores a 2.5 e podem ser outliers. Destes, somente o caso 5 possui o módulo de z-score maior que 3,0. Analisando cada caso individualmente, vimos que se trata de casos possíveis de serem encontrados na população e, dado o tamanho reduzido da amostra, podem dar a impressão de serem outliers. Por conta disso, os casos foram mantidos somente sob observação. No Apêndice B, a Tabela 42 contém os valores extremos de cada variável, retirada do SPSS®.

6.6. Análise de Fatores

Após o tratamento da amostra, iniciou-se a análise estatística propriamente dita. Na tentativa de identificar se a teoria dos constructos Poder, Experiência e Cultura pode ser aplicada no cenário brasileiro, executou-se uma análise de fatores com todas as variáveis, de modo que se conseguisse reduzir seu número e, possivelmente, explicá-las segundo a teoria da escala F-PEC.

Cinco fatores foram encontrados, a partir de uma análise exploratória cujo critério de corte dos fatores foi um *Eigenvalue* maior que 1. Porém, tanto a matriz de coeficientes quanto o *scree plot* sugerem que o número ideal de fatores seja quatro e não cinco, conforme a Tabela 6 e a Figura 5.

Component Matrix ^a					
	Component				
	1	2	3	4	5
Zscore: Poder 1	,129	-,015	,826	,209	-,100
Zscore: Poder 2	,476	-,106	,769	,201	,046
Zscore: Poder 3	,216	-,405	,719	-,180	,116
Zscore: Experiência 1	,184	,848	,135	-,023	,097
Zscore: Experiência 2	,236	,772	,063	-,022	,298
Zscore: Experiência 3	,338	,824	,108	-,040	,064
Zscore: Experiência 4	,187	,378	,014	,549	-,177
Zscore: Cultura 1	,567	-,162	-,143	,269	,466
Zscore: Cultura 2	,671	-,032	-,128	,392	-,368
Zscore: Cultura 3	,672	-,132	-,210	,521	-,180
Zscore: Cultura 4	,809	-,131	-,052	,043	,298
Zscore: Cultura 5	,817	-,144	-,226	,087	,306
Zscore: Cultura 6	,917	-,087	-,055	-,074	,117
Zscore: Cultura 7	,886	,135	,073	-,185	-,136
Zscore: Cultura 8	,805	,171	-,021	-,297	-,301
Zscore: Cultura 9	,828	,003	-,024	-,260	-,311
Zscore: Cultura 10	,913	-,176	-,099	-,110	,045
Zscore: Cultura 11	,874	,024	-,033	-,215	-,204
Zscore: Cultura 12	,885	-,183	-,035	-,068	,172

Extraction Method: Principal Component Analysis.

a. 5 components extracted.

Tabela 6: Matriz de coeficientes da análise de fatores.

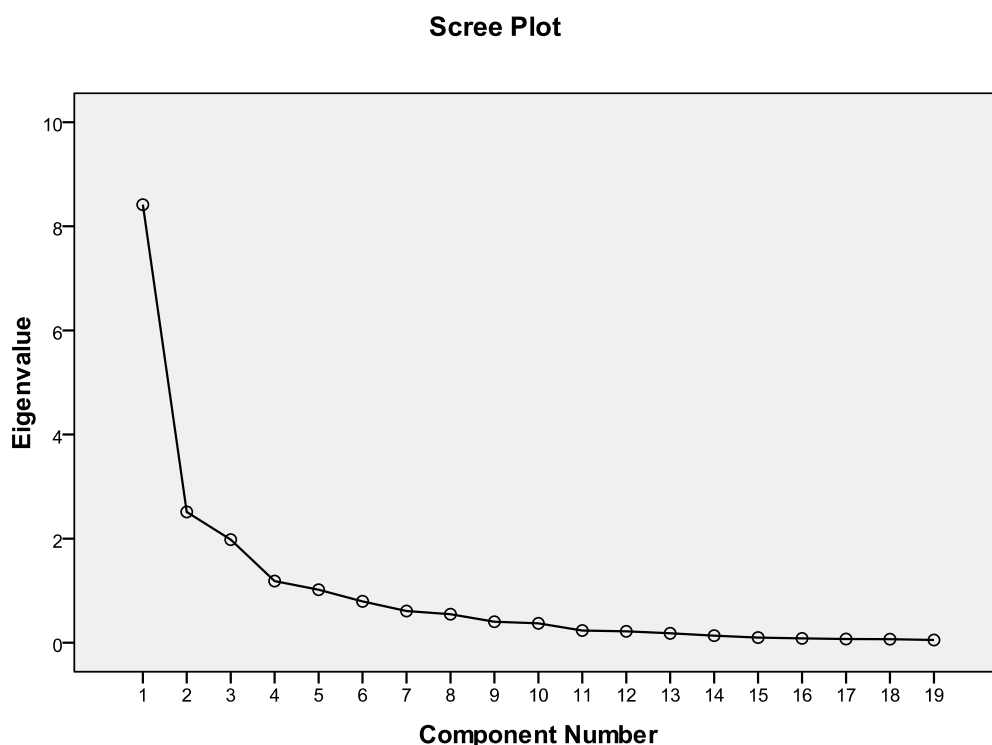


Figura 5: Gráfico *scree plot* da análise de fatores.

A Tabela 6 mostra que os maiores coeficientes de cada variável, marcados em negrito, formam somente quatro dos cinco fatores extraídos pelo procedimento. Pelo gráfico *scree plot*, pode-se observar que o “cotovelo”, ponto onde as diferenças entre Eigenvalue passam a ser menores, está na marca de quatro fatores. Pôde-se notar, também, que a variável E4 ficou isolada em um fator. Esta foi inserida no questionário, para medir a intenção de perpetuidade do controle familiar. Como a questão não constava do questionário original de Astrachan et al. (2002) e a variável resultante mostrou-se pouco correlacionada com as demais, ela foi excluída da análise, restando somente os três fatores que correspondem aos constructos Poder, Experiência e Cultura.

Diante disso, uma nova análise fatorial foi executada excluindo a variável E4, resultando na matriz a seguir, em que estão ressaltados os maiores coeficientes de cada variável. Desta vez, os coeficientes sugerem que a quantidade ideal de fatores é três. Os testes de Bartlett e Kaiser-Meyer-Olkin mostraram um nível de significância de ,000 e MSA maior que 80%, respectivamente. Isso nos permite concluir que o modelo da escala F-PEC tem uma adequação muito boa à

indústria brasileira e é possível usá-lo para identificar os principais tipos de empresa familiar encontrados no Brasil.

Component Matrix^a

	Component				
	1	2	3	4	5
Zscore: Poder 1	,127	-,036	,825	,162	-,174
Zscore: Poder 2	,476	-,119	,768	,216	-,045
Zscore: Poder 3	,222	-,392	,717	-,145	,188
Zscore: Experiência 1	,176	,852	,141	,084	,046
Zscore: Experiência 2	,231	,782	,069	,165	,225
Zscore: Experiência 3	,332	,840	,115	,088	-,007
Zscore: Cultura 1	,568	-,159	-,145	,455	,300
Zscore: Cultura 2	,670	-,045	-,130	,303	-,561
Zscore: Cultura 3	,670	-,159	-,213	,450	-,397
Zscore: Cultura 4	,811	-,119	-,053	,153	,259
Zscore: Cultura 5	,818	-,140	-,227	,174	,270
Zscore: Cultura 6	,917	-,076	-,055	-,051	,162
Zscore: Cultura 7	,884	,144	,075	-,244	-,034
Zscore: Cultura 8	,806	,199	-,018	-,354	-,196
Zscore: Cultura 9	,830	,026	-,023	-,343	-,205
Zscore: Cultura 10	,915	-,158	-,099	-,097	,093
Zscore: Cultura 11	,875	,041	-,032	-,281	-,103
Zscore: Cultura 12	,886	-,175	-,036	-,041	,223

Extraction Method: Principal Component Analysis.

a. 5 components extracted.

Tabela 7: Nova matriz de coeficientes da análise de fatores, excluindo-se a variável E4

KMO and Bartlett's Test

Kaiser-Meyer-Olkin Measure of Sampling Adequacy.	,826
Bartlett's Test of Sphericity	Approx. Chi-Square
	805,367
	df
	153
	Sig.
	,000

Tabela 8: Resultados dos testes de Bartlett e MSA.

Partindo deste ponto, para confirmar a formação dos constructos, foi feita uma análise fatorial com as variáveis de cada constructo separadamente. Cada resultado indicou que somente um fator é formado por constructo, reforçando a teoria. A seguir estão dispostos os resultados encontrados.

Constructo Poder

Component Matrix ^a		KMO and Bartlett's Test	
	Component		
	1		
Zscore: Poder 1	,809	Kaiser-Meyer-Olkin Measure of Sampling Adequacy.	,625
Zscore: Poder 2	,913	Bartlett's Test of Sphericity	Approx. Chi-Square
Zscore: Poder 3	,802		51,152
			df
			3
			Sig.
			,000

Extraction Method: Principal Component Analysis.

a. 1 components extracted.

Tabela 9: Matriz de componentes do constructo Poder e Resultados dos testes de Bartlett e MSA.

Constructo Experiência

Component Matrix ^a		KMO and Bartlett's Test	
	Component		
	1		
Zscore: Experiência 1	,891	Kaiser-Meyer-Olkin Measure of Sampling Adequacy.	,724
Zscore: Experiência 2	,856	Bartlett's Test of Sphericity	67,987
Zscore: Experiência 3	,903	Approx. Chi-Square	3
		df	,000
		Sig.	

Extraction Method: Principal Component Analysis.

a. 1 components extracted.

Tabela 10: Matriz de componentes do constructo Experiência e Resultados dos testes de Bartlett e MSA.

Constructo Cultura

Component Matrix ^a			KMO and Bartlett's Test	
	Component			
	1	2		
Zscore: Cultura 1	,580	,540	Kaiser-Meyer-Olkin Measure of Sampling Adequacy.	,724
Zscore: Cultura 2	,677	,154	Bartlett's Test of Sphericity	67,987
Zscore: Cultura 3	,686	,394	Approx. Chi-Square	3
Zscore: Cultura 4	,811	,229	df	,000
Zscore: Cultura 5	,838	,291	Sig.	
Zscore: Cultura 6	,920	,022		
Zscore: Cultura 7	,870	-,315		
Zscore: Cultura 8	,800	-,458		
Zscore: Cultura 9	,832	-,362		
Zscore: Cultura 10	,929	,011		
Zscore: Cultura 11	,877	-,309		
Zscore: Cultura 12	,895	,073		

Extraction Method: Principal Component Analysis.

a. 2 components extracted.

Tabela 11: Matriz de componentes do constructo Cultura e Resultados dos testes de Bartlett e MSA.

6.7. Cálculo dos Fatores

Para calcular os valores de cada fator, foi feita uma análise de dois métodos bastante usados: *Summated Scales* e *Factor Scores*. A tabela a seguir mostra os valores das variáveis FAC_P, FAC_E e FAC_C, formadas através do método de *Factor Scores* que leva em consideração o peso de cada variável para o fator.

ID	FAC_P	FAC_E	FAC_C	ID	FAC_P	FAC_E	FAC_C
1	0,47373	-1,46808	0,30203	27	0,61985	0,60556	0,49916
2	0,15987	-0,86872	-0,10716	28	-0,18585	-1,02444	-0,36959
3	0,07468	1,95474	1,06790	29	0,70781	-1,39230	0,04668
4	-0,23113	0,36458	-3,03899	30	1,29288	0,42610	0,18943
5	-3,06170	1,02050	0,48092	31	0,33373	0,24598	0,50913
6	-0,00965	0,26791	0,13034	32	0,51471	0,35142	1,02057
7	-2,46452	-0,67921	-0,39586	33	-1,71035	0,31104	0,00016
8	0,48361	-0,84646	-1,81819	34	1,04612	0,21747	0,57668
9	-1,01594	0,82786	0,16862	35	0,93984	0,39285	0,44216
10	1,00221	-0,26165	0,44148	36	-0,74456	0,64284	0,65345
11	0,02957	-0,35268	-2,86787	37	1,25258	0,11159	0,05310
12	-1,78034	-1,48847	-0,83895	38	1,30871	0,70853	0,50113
13	0,55584	0,59191	0,36460	39	1,37354	0,60139	0,47723
14	0,17397	-0,68238	0,08778	40	-0,63987	-0,19846	0,39991
15	-0,86739	2,41714	-0,07126	41	0,12815	-1,50252	0,28403
16	-0,48532	-1,24459	0,71680	42	0,11769	-1,31805	0,86386
17	-0,50977	1,99269	1,03278	43	0,47015	0,70876	0,88785
18	1,20907	2,33015	-2,79491	44	-0,39531	-1,26974	-0,10783
19	0,03523	0,50704	-3,06308	45	-1,32695	-0,85247	0,11691
20	-0,96312	0,93924	0,54119	46	-1,17113	0,25611	-0,42376
21	0,29455	0,89704	0,59676	47	0,67972	-1,05339	0,54529
22	-0,64701	0,20634	0,79666	48	0,74772	-0,77206	-0,18663
23	1,58345	0,09233	-0,24225	49	0,80600	-1,46601	0,60384
24	-0,52186	0,26692	0,63558	50	-0,71891	-1,13683	-0,05244
25	0,11006	-0,70976	0,07297	51	1,29008	-0,08068	-0,09772
26	-0,36447	0,41292	0,36950				

Tabela 12: Valores de cada variável latente formada pelos *factor scores*.

O método *Summated Scale* consiste na soma ou média dos valores das variáveis que compõem o fator. Para que se possa gerar uma *Summated Scale*, primeiro devem ser confirmadas, primeiramente, algumas premissas. A uni dimensionalidade foi confirmada na análise de fatores anterior, em que cada constructo continha somente um fator com todas as variáveis relacionadas a ele pela teoria. Em seguida, foi medido o Alfa de Cronbach para testar a confiabilidade, ou *reliability score*, de cada constructo. Conforme descrito na Tabela 13, todos obtiveram valores de alfa acima de 0,7, indicando alta confiabilidade.

Constructo Poder		Constructo Experiência		Constructo Cultura	
Reliability Statistics		Reliability Statistics		Reliability Statistics	
Cronbach's Alpha	N of Items	Cronbach's Alpha	N of Items	Cronbach's Alpha	N of Items
,794	3	,859	3	,951	12

Tabela 13: Confiabilidade dos constructos.

Em sequência, a validade, ou *validity*, foi testada usando o método da validade convergente, em que a escala é testada contra uma outra escala semelhante. Foi usada a escala *Factor Scores* para testar se a *Summated Scales* tem correlação com ela. A Tabela 14 mostra baixos níveis de significância entre os fatores, sugerindo uma ortogonalidade, e correlação significativa a 0,01 quando se comparam as variáveis latentes semelhantes de cada escala. Isso demonstra a validade da escala.

Correlations						
	SUM_P	SUM_E	SUM_C	FAC_P	FAC_E	FAC_C
Kendall's tau_b						
SUM_P Correlation Coefficient	1,000			,772**		
Sig. (2-tailed)	.			,000		
N	51			51		
SUM_E Correlation Coefficient	,028	1,000		,099	,763**	
Sig. (2-tailed)	,790	.		,332	,000	
N	51	51		51	51	
SUM_C Correlation Coefficient	,103	,236*	1,000	-,071	,118	,858**
Sig. (2-tailed)	,293	,021	.	,465	,223	,000
N	51	51	51	51	51	51
Spearman's rho						
SUM_P Correlation Coefficient	1,000			,912**		
Sig. (2-tailed)	.			,000		
N	51			51		
SUM_E Correlation Coefficient	,046	1,000		,155	,906**	
Sig. (2-tailed)	,750	.		,277	,000	
N	51	51		51	51	
SUM_C Correlation Coefficient	,157	,299*	1,000	-,099	,161	,966**
Sig. (2-tailed)	,272	,033	.	,490	,259	,000
N	51	51	51	51	51	51

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

* . Correlation is significant at the 0.05 level (2-tailed).

Tabela 14: Testes de correlação entre as escalas semelhantes.

Atendendo a todas as premissas, a *Summated Scales* é possível de ser usada e foi escolhida por sua maior facilidade na replicação em outras pesquisas. Três variáveis chamadas SUM_P, SUM_E e SUM_C foram então criadas a partir da média dos valores das variáveis que as compõem, resultando na Tabela 15.

ID	SUM_P	SUM_E	SUM_C	ID	SUM_P	SUM_E	SUM_C
1	0,78251	-1,07394	0,28584	27	0,39322	0,61224	0,26785
2	-0,03621	-1,07394	-0,05314	28	-0,16488	-1,07394	-0,22556
3	0,07344	2,01525	0,77988	29	0,74983	-1,07394	-0,01245
4	-1,10143	-0,51794	-2,47703	30	0,97870	0,61224	0,02970
5	-2,46687	0,61224	0,56547	31	0,40432	0,61224	0,35795
6	-0,06830	0,03306	0,09317	32	0,66376	0,61224	0,77988
7	-2,03130	-1,07394	-0,13352	33	-1,43225	0,03306	0,08044
8	0,18732	-1,07394	-1,42720	34	0,97870	0,61224	0,35955
9	-0,80559	0,61224	0,22182	35	0,74983	0,61224	0,25081
10	0,97870	0,03306	0,29250	36	-0,50914	0,61224	0,56845
11	-0,67939	-1,07394	-2,21084	37	0,97870	0,61224	-0,10308
12	-1,20935	-1,07394	-0,52115	38	0,97870	0,88774	0,25133
13	0,46844	0,61224	0,23756	39	0,97870	0,61224	0,23406
14	0,22381	-0,52294	0,05192	40	-0,28558	0,06125	0,40174
15	-1,17297	1,90495	-0,13415	41	0,47577	-1,07394	0,30572
16	-0,00570	-1,07394	0,67340	42	0,55348	-1,07394	0,77988
17	-0,42064	1,87688	0,77988	43	0,55348	0,89024	0,66528
18	-0,21512	1,45533	-2,47703	44	-0,02614	-1,07394	0,06741
19	-0,96836	-0,49475	-2,47703	45	-0,69684	-0,51794	0,23877
20	-0,77669	0,89024	0,46437	46	-1,03920	0,06125	-0,24439
21	0,29835	0,89024	0,39893	47	0,97870	-0,51794	0,45579
22	-0,34209	0,05624	0,72460	48	0,66781	-0,51794	-0,17147
23	0,97870	0,06125	-0,37643	49	0,97870	-1,07394	0,51313
24	-0,37005	0,03306	0,52055	50	-0,30795	-1,07394	0,07968
25	0,35692	-0,51794	0,12499	51	0,97870	0,05624	-0,19126
26	-0,25729	0,05624	0,33345				

Tabela 15: Valores das variáveis latentes gerados por *summated scales*.

6.8. Análise de *Clusters*

A fim de identificar os principais tipos encontrados na amostra, foi feita uma análise de *cluster* para agrupar os casos de maior semelhança. Inicialmente foi empregada a técnica hierárquica que sugeriu uma quantidade entre quatro e seis grupos, posteriormente confirmados na técnica não hierárquica de *k-means*.

Antes de aplicar o agrupamento, uma nova análise em busca de *outliers* se fez necessária, devido à natureza multivariada da análise de *clusters*. Por causa disso, a distância D^2 de *Mahalanobis* foi calculada para cada caso. Os maiores valores encontrados foram dos casos 5, com $D^2=10,956$, e 18, com $D^2=15,103$, ambos inferiores ao valor crítico de 16,266 para 3 graus de liberdade. A tabela 14 mostra os valores encontrados.

Para testar se existe multicolinearidade entre as variáveis usadas na análise de *cluster*, foi feito um teste de Regressão Múltipla com estatísticas de colinearidade. O resultado indicou que não existe colinearidade entre as variáveis, pois foram encontrados valores de tolerância e VIF muito próximos de 1, conforme indicados na Tabela 16.

ID	D^2	ID	D^2	ID	D^2
1	2,52163	18	15,10259	35	1,26063
2	1,51820	19	9,50531	36	1,42166
3	5,46660	20	2,29063	37	2,15274
4	9,66045	21	1,17763	38	2,36710
5	10,95569	22	1,24392	39	1,83244
6	0,02618	23	2,00787	40	0,47704
7	7,73380	24	0,81099	41	2,07367
8	4,28874	25	0,56114	42	3,14528
9	1,56699	26	0,34856	43	1,66277
10	1,35715	27	0,69661	44	1,60279
11	7,92116	28	1,51031	45	1,44042
12	3,49906	29	2,29971	46	1,53145
13	0,78035	30	1,98055	47	1,88545
14	0,44051	31	0,74015	48	1,07896
15	6,74196	32	1,52817	49	3,25441
16	2,76923	33	3,19065	50	1,80877
17	5,27594	34	1,81173	51	1,67424

Tabela 16: Distâncias quadradas de *Mahalanobis*.

6.9. Método Hierárquico

Para identificar a quantidade de tipos encontrados na amostra, foi executada uma análise hierárquica de clusters, como forma exploratória. Foram usados dois métodos hierárquicos *Average Linkage* para identificar a quantidade ideal de clusters. Primeiramente, foi executada uma análise usando o método *Between-groups* e em seguida o método *Within Groups*. Ambos usaram a distância Euclidiana quadrada para formar os grupos. Pôde ser observado, conforme os resultados do SPSS® dispostos a seguir na Tabela 17 e Tabela 18, que, na primeira análise, a maior variação de heterogeneidade ocorreu do passo 45 para o passo 46, sugerindo 6 clusters. A segunda análise mostrou maior variação da heterogeneidade do passo 45 para o passo 48, sugerindo 4 clusters. Podemos então concluir que a quantidade ideal de grupos está entre 4 e 6.

Estágio	Coeficiente	% Mudança	Grupos
40	1,056	21%	11
41	1,179	12%	10
42	1,902	61%	9
43	2,404	26%	8
44	2,610	9%	7
45	2,815	8%	6
46	4,563	62%	5
47	5,849	28%	4
48	6,072	4%	3
49	7,212	19%	2
50	9,366	30%	1

Tabela 17: Coeficientes de aglomeração dos últimos estágios da análise hierárquica de clusters usando o método *Between-groups*.

Estágio	Coeficiente	% Mudança	Grupos
40	,764	10%	11
41	,826	8%	10
42	1,010	22%	9
43	1,123	11%	8
44	1,281	14%	7
45	1,710	34%	6
46	1,846	8%	5
47	2,054	11%	4
48	2,915	42%	3
49	3,493	20%	2
50	4,292	23%	1

Tabela 18: Coeficientes de aglomeração dos últimos estágios da análise hierárquica de clusters usando o método *Within groups*.

6.10. Método *K-Means*

6.10.1. Seis Grupos

O método não hierárquico *K-Means Cluster* foi executado para seis, cinco e quatro grupos. Para seis grupos, o procedimento resultou em dois grupos com somente um caso. A menor distância entre centros de grupos foi de 1,466, entre os grupos 3 e 6. A maior distância foi de 3,878, entre os grupos 4 e 5. A Tabela 19, Tabela 20 e Tabela 21, extraídas do SPSS®, demonstram o resultado encontrado.

Distances between Final Cluster Centers

Cluster	1	2	3	4	5	6
1		2,149	3,151	2,309	3,559	2,555
2	2,149		2,146	3,255	1,771	1,627
3	3,151	2,146		2,928	2,757	1,466
4	2,309	3,255	2,928		3,878	3,405
5	3,559	1,771	2,757	3,878		3,037
6	2,555	1,627	1,466	3,405	3,037	

Tabela 19: Distâncias entre os seis centroides.

Number of Cases in each Cluster

Cluster	1	4,000
	2	5,000
	3	18,000
	4	1,000
	5	1,000
	6	22,000
Valid		51,000
Missing		,000

Tabela 20: Casos por grupo.

Final Cluster Centers

	Cluster					
	1	2	3	4	5	6
SUM_P	-,64046	-1,28179	,26741	-,21512	-2,46687	,31088
SUM_E	-,79014	-,51430	,89390	1,45533	,61224	-,56481
SUM_C	-2,14803	-,11597	,35611	-2,47703	,56547	,21243

Tabela 21: Posição dos seis centroides.

Analisando os centroides finais dos grupos, à luz da taxonomia proposta, pôde ser observado que os grupos 2 e 5 pertenceriam à mesma classe denominada Classe 5, enquanto que os grupos 3 e 6 pertenceriam à mesma classe denominada Classe 14.

A Análise Multivariada da Variância (MANOVA) indicou que os grupos são estatisticamente diferentes para as quatro medidas de diferenças usadas no teste, *Pillai's Trace*, *Wilk's Lambda*, *Hotelling's Trace* e *Roy's Largest Root*. Foi

adotado o *Wilk's Lambda* por ser o valor mais comumente usado nesse tipo de comparação. O valor do lambda de Wilk ficou em 0,019, conforme na Tabela 22.

<i>Multivariate Tests^c</i>						
<i>Effect</i>		<i>Value</i>	<i>F</i>	<i>Hypothesis df</i>	<i>Error df</i>	<i>Sig.</i>
<i>Intercept</i>	<i>Pillai's Trace</i>	,604	21,862 ^a	3,000	43,000	,000
	<i>Wilks' Lambda</i>	,396	21,862 ^a	3,000	43,000	,000
	<i>Hotelling's Trace</i>	1,525	21,862 ^a	3,000	43,000	,000
	<i>Roy's Largest Root</i>	1,525	21,862 ^a	3,000	43,000	,000
<i>QCL_1</i>	<i>Pillai's Trace</i>	2,070	20,043	15,000	135,000	,000
	<i>Wilks' Lambda</i>	,019	25,432	15,000	119,105	,000
	<i>Hotelling's Trace</i>	10,145	28,180	15,000	125,000	,000
	<i>Roy's Largest Root</i>	6,643	59,786 ^b	5,000	45,000	,000

a. Exact statistic

b. The statistic is an upper bound on F that yields a lower bound on the significance level.

c. Design: Intercept + QCL_1

Tabela 22: Teste de MANOVA.

6.10.2. Cinco Grupos

Continuando a análise, foi executado o procedimento *K-Means Cluster* para cinco grupos, que resultou na Tabela 23, Tabela 24 e Tabela 25. Pôde-se observar ainda um grupo com somente um caso. A menor distância entre centroides foi de 1,628, entre os grupos 3 e 1. A maior distância entre os centros foi de 3,307, entre os grupos 3 e 2.

Analisando os centroides finais à luz da taxonomia proposta, pôde-se observar que cada grupo pertence a somente uma classe. Os grupos 1 a 5 correspondem às classes 14, 13, 11, 10 e 5, respectivamente.

Os resultados da MANOVA indicam que os grupos são estatisticamente diferentes com um valor de *Wilk's Lambda* de 0,032, conforme indicado na Tabela 26.

Distances between Final Cluster Centers

Cluster	1	2	3	4	5
1		2,684	1,628	3,236	1,948
2	2,684		3,307	2,309	2,350
3	1,628	3,307		2,958	1,968
4	3,236	2,309	2,958		3,301
5	1,948	2,350	1,968	3,301	

Tabela 23: Distâncias entre os cinco centroides.

Number of Cases in each Cluster

Cluster	1	32,000
	2	4,000
	3	8,000
	4	1,000
	5	6,000
Valid		51,000
Missing		,000

Tabela 24: Casos por grupo.

Final Cluster Centers

	Cluster				
	1	2	3	4	5
SUM_P	,45039	-,64046	-,34497	-,21512	-1,47930
SUM_E	-,18837	-,79014	1,21154	1,45533	-,32654
SUM_C	,22935	-2,14803	,46806	-2,47703	-,00240

Tabela 25: Posição dos cinco centroides.

Multivariate Tests^c

Effect		Value	F	Hypothesis df	Error df	Sig.
Intercept	Pillai's Trace	,705	35,066 ^a	3,000	44,000	,000
	Wilks' Lambda	,295	35,066 ^a	3,000	44,000	,000
	Hotelling's Trace	2,391	35,066 ^a	3,000	44,000	,000
	Roy's Largest Root	2,391	35,066 ^a	3,000	44,000	,000
QCL_3	Pillai's Trace	1,895	19,722	12,000	138,000	,000
	Wilks' Lambda	,032	26,130	12,000	116,705	,000
	Hotelling's Trace	8,310	29,548	12,000	128,000	,000
	Roy's Largest Root	5,940	68,304 ^b	4,000	46,000	,000

a. Exact statistic

b. The statistic is an upper bound on F that yields a lower bound on the significance level.

c. Design: Intercept + QCL_3

Tabela 26: Teste de MANOVA.

6.10.3. Quatro Grupos

A análise com quatro *clusters* não mostrou grupos com menos de cinco casos. A menor distância entre os centros finais dos grupos foi de 1,483, entre os *clusters* 3 e 1. A maior distância entre os centroides foi de 2,999, entre os grupos 4 e 3. A Tabela 27, Tabela 28 e Tabela 29 descrevem o resultado.

Analisando os centroides finais à luz da taxonomia proposta, pôde-se observar que os grupos 1 e 3 pertencem à mesma classe 14. Os resultados da MANOVA indicaram que os grupos são diferentes estatisticamente entre si, com um lambda de Wilk igual a 0,033, conforme a Tabela 30.

Distances between Final Cluster Centers

Cluster	1	2	3	4
1		1,744	1,483	2,586
2	1,744		2,065	2,396
3	1,483	2,065		2,999
4	2,586	2,396	2,999	

Tabela 27: Distâncias entre os quatro centroides.**Number of Cases in each Cluster**

Cluster	1	22,000
	2	7,000
	3	17,000
	4	5,000
Valid		51,000
Missing		,000

Tabela 28: Casos por grupo.**Final Cluster Centers**

	Cluster			
	1	2	3	4
SUM_P	,31088	-1,38306	,33052	-,55540
SUM_E	-,56481	-,19243	,91047	-,34105
SUM_C	,21243	,02963	,36401	-2,21383

Tabela 29: Posição dos quatro centroides.

Multivariate Tests^c						
<i>Effect</i>		<i>Value</i>	<i>F</i>	<i>Hypothesis df</i>	<i>Error df</i>	<i>Sig.</i>
<i>Intercept</i>	<i>Pillai's Trace</i>	,598	22,326 ^a	3,000	45,000	,000
	<i>Wilks' Lambda</i>	,402	22,326 ^a	3,000	45,000	,000
	<i>Hotelling's Trace</i>	1,488	22,326 ^a	3,000	45,000	,000
	<i>Roy's Largest Root</i>	1,488	22,326 ^a	3,000	45,000	,000
<i>QCL_5</i>	<i>Pillai's Trace</i>	1,874	26,079	9,000	141,000	,000
	<i>Wilks' Lambda</i>	,033	37,180	9,000	109,669	,000
	<i>Hotelling's Trace</i>	8,293	40,238	9,000	131,000	,000
	<i>Roy's Largest Root</i>	6,132	96,074 ^b	3,000	47,000	,000

a. Exact statistic

b. The statistic is an upper bound on F that yields a lower bound on the significance level.

c. Design: Intercept + QCL_5

Tabela 30: Teste de MANOVA.

6.10.4. Quantidade Ideal

Para decidir qual quantidade de grupos é a ideal, analisamos três características de cada procedimento *K-Means* executado anteriormente. Primeiro, buscou-se observar o valor de diferenciação dos *clusters*, expresso pelo lambda de *Wilk*. Os valores encontrados, conforme as tabelas anteriores, foram 0,019, 0,032 e 0,033, para as análises com seis, cinco e quatro grupos, respectivamente. Pôde-se notar uma diferença significativa entre os valores de 0,019 e 0,033, sugerindo que a diferenciação entre os *clusters* aumentou sensivelmente de seis para cinco grupos. Já de cinco para quatro, o valor de lambda de *Wilk* passou de 0,033 para 0,034, sugerindo uma alteração menor nessa diferença entre grupos. Entretanto, quando observada a estatística F, pôde-se observar uma variação maior de cinco para quatro grupos (26,130 para 37,180) do que de seis para cinco grupos (25,432 para 26,130).

Em seguida, foram observadas as distâncias entre os centroides finais de cada análise. Como o objetivo é encontrar *clusters* com maior distância intergrupos, buscou-se a maior das menores distâncias entre os centroides. Desse modo, percebeu-se que a configuração com cinco grupos teve a maior das menores distâncias entre grupos, indicando que os dois grupos mais próximos dessa configuração estão mais afastados entre si que os grupos mais próximos das demais configurações. Os valores para as menores distâncias foram de 1,466, 1,628 e 1,483, para as configurações com seis, cinco e quatro grupos, respectivamente.

Por fim, foram analisadas as estatísticas descritivas das distâncias de cada caso ao centro do grupo a que pertence. Como outro objetivo da análise de *clusters* é identificar grupos com a menor distância entre seus membros, buscou-se observar qual configuração resultou na menor das maiores distâncias entre membros de cada grupo. As tabelas a seguir, extraídas do SPSS®, mostram as estatísticas descritivas para essas distâncias, agrupadas por *cluster membership*.

Para a configuração com seis grupos, pôde-se observar que a maior média entre as distâncias de membros do mesmo grupo a seu centroide foi de 0,778, enquanto que para cinco grupos foi de 0,874 e para quatro grupos foi de 1,008. Comparando as maiores distâncias entre um membro do grupo e seu centroide,

para cada configuração, pôde ser observado que a configuração com cinco grupos possuía a menor entre essas distâncias. O caso mais afastado do seu centroide para essa configuração possuía uma distância de 1,476, enquanto que para quatro grupos a distância era de 1,847 e, para seis grupos, era de 1,827.

Seis grupos

Descriptives ^{a,b}						
Cluster Number of Case			Statistic	Std. Error		
Distance of Case from its Classification Cluster Center	1	Mean	,6514546	,17598108		
		95% Confidence Interval for Mean	,0914043			
		Lower Bound	1,2115050			
		Upper Bound	,6445609			
		5% Trimmed Mean	,5894111			
		Median	,124			
		Variance	,35196217			
		Std. Deviation	,29326			
		Minimum	1,13374			
		Maximum	,84048			
		Range	,64983			
		Interquartile Range	,997	1,014		
		Skewness	1,850	2,619		
		Kurtosis				
	2	Mean	,7105417	,05872066		
		95% Confidence Interval for Mean	,5475070			
		Lower Bound	,8735764			
		Upper Bound	,7041442			
		5% Trimmed Mean	,6841174			
		Median	,017			
		Variance	,13130340			
		Std. Deviation	,60068			
		Minimum	,93556			
		Maximum	,33487			
		Range	,19596			
		Interquartile Range	1,777	,913		
		Skewness	3,514	2,000		
		Kurtosis				
	3	Mean	,7779946	,09997915		
		95% Confidence Interval for Mean	,5670571			
		Lower Bound	,9889322			
		Upper Bound	,7600069			
		5% Trimmed Mean	,7698721			
		Median	,180			
		Variance	,42417560			
		Std. Deviation	,05295			
		Minimum	1,82682			
		Maximum	1,77386			
		Range	,65927			
		Interquartile Range	,620	,536		
		Skewness	,890	1,038		
		Kurtosis				
	6	Mean	,7409195	,05158738		
		95% Confidence Interval for Mean	,6336377			
		Lower Bound	,8482014			
		Upper Bound	,7565679			
		5% Trimmed Mean	,7781524			
		Median	,059			
		Variance	,24196628			
		Std. Deviation	,10937			
		Minimum	1,08843			
		Maximum	,97906			
		Range	,23383			
		Interquartile Range	-1,233	,491		
		Skewness	1,834	,953		
		Kurtosis				

a. Distance of Case from its Classification Cluster Center is constant when Cluster Number of Case = 4. It has been omitted.

b. Distance of Case from its Classification Cluster Center is constant when Cluster Number of Case = 5. It has been omitted.

Tabela 31: Análise descritiva das distâncias até o centro para seis grupos.

Extreme Values ^{a,b,c}				
Cluster Number of Case	Case Number	ID	Value	
1 Highest 1	8	8	1,13374	
2 Highest 1	7	7	,93556	
3 Highest 1	15	15	1,82682	
6 Highest 1	23	23	1,08843	

a. The requested number of extreme values exceeds the number of data points. A smaller number of extremes is displayed.

b. Distance of Case from its Classification Cluster Center is constant when Cluster Number of Case = 4. It has been omitted.

c. Distance of Case from its Classification Cluster Center is constant when Cluster Number of Case = 5. It has been omitted.

Tabela 32: Maiores distâncias encontradas.

Cinco grupos

Descriptives ^a				Statistic	Std. Error
Cluster Number of Case					
Distance of Case from its Classification Cluster Center	1	Mean		,8740250	,03710762
		95% Confidence Interval for Mean	Lower Bound	,7983435	
			Upper Bound	,9497065	
		5% Trimmed Mean		,8833896	
		Median		,9228972	
		Variance		,044	
		Std. Deviation		,20991241	
		Minimum		,35811	
		Maximum		1,19900	
		Range		,84089	
		Interquartile Range		,25176	
		Skewness		-,701	,414
		Kurtosis		,066	,809
	2	Mean		,6514546	,17598108
		95% Confidence Interval for Mean	Lower Bound	,0914043	
			Upper Bound	1,2115050	
		5% Trimmed Mean		,6445609	
		Median		,5894111	
		Variance		,124	
		Std. Deviation		,35196217	
		Minimum		,29326	
		Maximum		1,13374	
		Range		,84048	
		Interquartile Range		,64983	
		Skewness		,997	1,014
		Kurtosis		1,850	2,619
	3	Mean		,8240975	,07892431
		95% Confidence Interval for Mean	Lower Bound	,6374712	
			Upper Bound	1,0107239	
		5% Trimmed Mean		,8170686	
		Median		,7668154	
		Variance		,050	
		Std. Deviation		,22323167	
		Minimum		,53817	
		Maximum		1,23655	
		Range		,69838	
		Interquartile Range		,31765	
		Skewness		,745	,752
		Kurtosis		,354	1,481
	5	Mean		,8684838	,15084550
		95% Confidence Interval for Mean	Lower Bound	,4807231	
			Upper Bound	1,2562445	
		5% Trimmed Mean		,8623053	
		Median		,8895997	
		Variance		,137	
		Std. Deviation		,36949451	
		Minimum		,37201	
		Maximum		1,47617	
		Range		1,10417	
		Interquartile Range		,51188	
		Skewness		,536	,845
		Kurtosis		1,268	1,741

a. Distance of Case from its Classification Cluster Center is constant when Cluster Number of Case = 4. It has been omitted.

Tabela 33: Análise descritiva das distâncias até o centro para cinco grupos.

Extreme Values ^{a,b}				
Cluster Number of Case		Case Number	ID	Value
Distance of Case from its Classification Cluster Center	1	Highest 1	38	1,19900
	2	Highest 1	8	1,13374
	3	Highest 1	15	1,23655
	5	Highest 1	5	1,47617

a. The requested number of extreme values exceeds the number of data points. A smaller number of extremes is displayed.

b. Distance of Case from its Classification Cluster Center is constant when Cluster Number of Case = 4. It has been omitted.

Tabela 34: Maiores distâncias encontradas.

Quatro grupos

Descriptives						
Cluster Number of Case			Statistic	Std. Error		
Distance of Case from its Classification Cluster Center	1	Mean	,7409195	,05158738		
		95% Confidence Interval for Mean	,6336377			
		Lower Bound	,8482014			
		Upper Bound	,7565679			
		5% Trimmed Mean	,7781524			
		Median	,059			
		Variance	,24196628			
		Std. Deviation	,10937			
		Minimum	1,08843			
		Maximum	,97906			
		Range	,23383			
		Interquartile Range	-1,233	,491		
		Skewness	1,834	,953		
		Kurtosis				
	2	Mean	,8790149	,15330492		
		95% Confidence Interval for Mean	,5038912			
		Lower Bound	1,2541385			
		Upper Bound	,8828690			
		5% Trimmed Mean	1,0089126			
		Median	,165			
		Variance	,40560670			
		Std. Deviation	,23632			
		Minimum	1,45233			
		Maximum	1,21601			
		Range	,59867			
		Interquartile Range	-,384	,794		
		Skewness	-,199	1,587		
		Kurtosis				
	3	Mean	,7457835	,10806218		
		95% Confidence Interval for Mean	,5167019			
		Lower Bound	,9748651			
		Upper Bound	,7218813			
		5% Trimmed Mean	,7135073			
		Median	,199			
		Variance	,44555178			
		Std. Deviation	,05161			
		Minimum	1,87020			
		Maximum	1,81859			
		Range	,64937			
		Interquartile Range	,895	,550		
		Skewness	1,197	1,063		
		Kurtosis				
	4	Mean	1,0083843	,25000646		
		95% Confidence Interval for Mean	,3142550			
		Lower Bound	1,7025135			
		Upper Bound	,9892917			
		5% Trimmed Mean	,7433118			
		Median	,313			
		Variance	,55903145			
		Std. Deviation	,51326			
		Minimum	1,84717			
		Maximum	1,33391			
		Range	1,00460			
		Interquartile Range	1,009	,913		
		Skewness	-,473	2,000		
		Kurtosis				

Tabela 35: Análise descritiva das distâncias até o centro para quatro grupos.

Extreme Values ^a					
Cluster Number of Case			Case Number	ID	Value
Distance of Case from its Classification Cluster Center	1	Highest	1	23	1,08843
	2	Highest	1	5	1,45233
	3	Highest	1	15	1,87020
	4	Highest	1	18	1,84717

a. The requested number of extreme values exceeds the number of data points. A smaller number of extremes is displayed.

Tabela 36: Maiores distâncias encontradas.

A partir dessas três características e buscando manter-se a parcimônia na escolha do modelo ideal, nota-se que a configuração com cinco grupos tem maior ganho de lambda de *Wilk*, apesar do maior ganho em *F* ser para quatro grupos. O modelo de cinco grupos também é o que possui a maior distância entre os grupos mais próximos, sugerindo que seus grupos estão mais afastados entre si. Por último, apesar de não possuir a menor das maiores médias de distâncias intragrupos, possui a menor das maiores distâncias que um elemento possui até o seu centroide, sugerindo que os casos estão mais próximos dos seus centroides que nos demais modelos.

O modelo escolhido de cinco grupos possui um grupo com somente um caso, o que pode não parecer ideal. Entretanto, devido ao tamanho reduzido da amostra, 1 caso pode ser representativo de quase 2% da população e, pensando nisso, consideramos como uma opção válida.

6.11. Interpretação dos Resultados

Os cinco grupos encontrados correspondem às classes 14, 13, 11, 10 e 5. A primeira é a classe de maior concentração e que denota a configuração de Poder, Experiência e Cultura mais comumente encontrada na população. Com sua característica mediana nas três subescalas, pode-se definir como uma empresa comum, ou corriqueira no setor.

A Classe 13 possui a subescala cultura mais baixa que a anterior. Esta subescala mede a sinergia entre os valores da família e da empresa, bem como o comprometimento da primeira com a segunda. Pode-se imaginar uma empresa onde exista desavença de pensamentos entre os membros familiares que a controlam. Trata-se do segundo grupo mais concentrado.

A Classe 11 possui Experiência inferior à classe de comparação 14. A ausência de sucessão e a centralização do poder de decisão em uma geração somente reduz o valor da subescala Experiência. Por causa disso, podem-se imaginar os membros desta Classe como empresas jovens, ou que tenham somente um dono.

A classe seguinte, Classe 10, possui tanto Experiência, quanto Cultura abaixo da classe de comparação, Classe 14. Pode-se imaginar uma mistura das duas classes, 13 e 11, onde empresas com controle centralizado em poucos membros da mesma geração, que possuem discórdia de pensamentos e comprometimento. Esta classe foi a que se encontrou somente um elemento na amostra, entretanto não é difícil de imaginar uma empresa dirigida por primos ou irmãos, que pouco concordam na maneira de guiá-la.

Por fim, a Classe 5 possui o valor da subescala Poder abaixo da Classe 14. Esta subescala refere-se ao poder da família exercido através da propriedade. Podem-se imaginar, nesta classe, empresas em processo de descentralização, ou até mesmo IPO, onde buscam diversificar a propriedade como forma de evolução.

6.12. Comparação com Outras Variáveis

Na pesquisa enviada aos presidentes de empresas, foram incluídas questões quanto ao ano de fundação e à quantidade aproximada de funcionários. Uma análise dessas variáveis para identificar uma possível correlação com o grupo a que a empresa pertence foi feita e sugere que exista essa relação para a quantidade de funcionários.

Foi feita uma análise *crosstabs* para identificar se existe correlação entre as variáveis *cluster membership* e ano de fundação, bem como entre *cluster membership* e quantidade de funcionários. As tabelas a seguir demonstram o resultado encontrado para o coeficiente Eta. Este é usado para medir a correlação entre uma variável categórica (*cluster membership*), definida como independente, e uma variável numérica (ano de fundação ou quantidade de funcionários), definida como dependente. Valores próximos de 1 indicam uma possível correlação e, nos casos analisados, foram encontrados valores de 0,404 e 0,504 para as correlações com o ano de fundação e quantidade de funcionários, respectivamente. A Tabela 37 e a Tabela 38 mostram os coeficientes encontrados.

Como o coeficiente Eta possui valores mais próximos de 1 do que de 0 para a variável quantidade de funcionários, sugere que possa existir alguma correlação entre elas, mesmo que fraca. Já para o ano de fundação, o valor do coeficiente está mais próximo de 0 do que de 1, sugerindo a ausência dessa relação.

Directional Measures			Value
Nominal by Interval	Eta	Ano de Fundação Dependent	,404
		Cluster Number of Case Dependent	,821

Tabela 37: Coeficiente Eta para o Ano de Fundação e *cluster membership*.

Directional Measures			Value
Nominal by Interval	Eta	Quantidade de Funcionários Dependent	,540
		Cluster Number of Case Dependent	,870

Tabela 38: Coeficiente Eta para o Quantidade de Funcionários e *cluster membership*.